

MASTERARBEIT

IMPLEMENTIERUNG VON KÜNSTLICHER INTELLIGENZ IM BEARBEITUNGSPROZESS VON SERVICEANFRAGEN IM KUNDENSUPPORT

am Beispiel der Styria IT Solutions GmbH & Co KG

ausgeführt am



Studiengang

Informationstechnologien und Wirtschaftsinformatik

Von: Bernhard Glanznig

Personenkennzeichen: 2210320007

Graz, am 18. März 2024

Glanznig Bernhard
.....
Unterschrift

EHRENWÖRTLICHE ERKLÄRUNG

Ich erkläre ehrenwörtlich, dass ich die vorliegende Arbeit selbstständig und ohne fremde Hilfe verfasst, andere als die angegebenen Quellen nicht benützt und die benutzten Quellen wörtlich zitiert sowie inhaltlich entnommene Stellen als solche kenntlich gemacht habe.

Glanznig Bernhard
.....
Unterschrift

DANKSAGUNG

Mein Dank gilt allen Personen, die mich auf meinem Weg zum Abschluss die letzten 5 Jahre begleitet haben. Speziell danken möchte ich meiner Frau, meiner Familie und meinen Freunden, die mich immer unterstützt haben, mir halt geben und immer hinter mir gestanden sind.

Ein besonderer Dank gebührt auch der Firma Styria IT Solutions GmbH und den Expert*innen, die mich bei meiner Forschung unterstützt haben. Die Möglichkeit, mit Ihnen zusammenzuarbeiten und von Ihrem Wissen zu profitieren, war von unschätzbarem Wert für den Erfolg dieser Arbeit.

Vielen Dank auch an Herrn Mag. Berndt Jesenko für die Unterstützung während der Masterarbeit und der Unterstützung in der Strukturierung der vorliegenden Arbeit.

Danke auch allen Anderen, die in irgendeiner Weise zu dieser Arbeit beigetragen haben, sei es durch Diskussionen, Feedback oder auf andere Weise.

KURZFASSUNG

Die vorliegende Masterarbeit widmet sich der Implementierung künstlicher Intelligenz (KI) im Bearbeitungsprozess von Serviceanfragen im Kundensupport. Dabei wird der Design Science Research Ansatz als Forschungsmethode verwendet, um eine effektive Gestaltung und Umsetzung der KI-Integration zu gewährleisten.

Zu Beginn der Arbeit wird der Design Science Research Ansatz eingehend erläutert. Dieser methodische Rahmen ermöglicht die Entwicklung und Evaluation von artefaktbasierten Lösungen, in diesem Fall die Implementierung von KI im Kundensupport. Es werden Designziele definiert und eine iterative Vorgehensweise skizziert, die die praxisnahe Anwendung von Forschungsergebnissen gewährleistet.

Im Anschluss bezieht sich der Fokus auf das Incident Management im Kundensupport. Es erfolgt eine detaillierte Analyse des bestehenden Prozesses, um Herausforderungen zu identifizieren, die durch KI gelöst werden können. Der Schwerpunkt liegt auf der Automatisierung von Routinetätigkeiten, der schnellen Klassifizierung von Anfragen und der Verbesserung der Wissensbasis.

Abgeschlossen werden die Grundlagen mit dem Thema künstliche Intelligenz. Verschiedene KI-Technologien wie Supervised Learning und der KNN-Algorithmus werden erläutert, um eine fundierte Basis für die spätere Implementierung zu schaffen.

Das vierte Kapitel widmet sich einer empirischen Studie zur Evaluation der implementierten KI im Kundensupport. Hierbei wird die ausgewählte KI-Technologie in der Praxis getestet, und die Ergebnisse werden auf ihre Wirksamkeit und Effizienz hin analysiert. Die Erkenntnisse aus der empirischen Studie fließen in die abschließende Reflexion und Handlungsempfehlungen für zukünftige Implementierungen im Kundensupport ein. Eine Schlussfolgerung bzw. der Ausblick schließen diese Arbeit ab.

ABSTRACT

This master's thesis investigates implementing artificial intelligence (AI) to process service requests in customer support. The Design Science Research approach ensures effective design and implementation of AI integration and is explained in detail. Design goals are defined, and an iterative approach outlined that ensures the practical application of research results. The focus then turns to incident management in customer support.

A detailed analysis of the existing process is carried out to identify challenges that AI can solve. The focus is on the automation of routine activities, the rapid classification of queries, and the improvement of the knowledge base. The third chapter covers the basics of AI. Various AI technologies, such as supervised learning and the KNN algorithm, are explained to create a foundation for subsequent implementation. The fourth chapter is dedicated to an empirical study to evaluate the AI implemented in customer support. The selected AI technology is tested in practice, and the results are analyzed in terms of their effectiveness and efficiency. The findings from the empirical study are incorporated into the concluding reflection and recommendations for action for future implementations in customer support.

INHALTSVERZEICHNIS

1	EINLEITUNG	1
1.1	Unternehmenspräsentation	1
1.2	Ausgangssituation der Masterarbeit	2
1.3	Zielsetzung der Masterarbeit	3
1.4	Vorgehensweise	3
2	METHODOLOGIE	5
2.1	Design Science Research Ansatz	5
2.1.1	Begriffsdefinition	6
2.1.2	Gestaltungsorientierte Wirtschaftsinformatik	6
2.1.3	Prozess in der gestaltungsorientierten Wirtschaftsinformatik	7
3	THEORIE	10
3.1	Incident Management	10
3.1.1	Begriffsdefinition	10
3.1.2	Zusammenhang von Incident Management mit ITSM bzw. ITIL	10
3.1.3	Herausforderungen und Ziele des Incident Managements	14
3.2	Künstliche Intelligenz	16
3.2.1	Geschichte der künstlichen Intelligenz	16
3.2.2	Big Data	18
3.2.3	Praktische Anwendungsgebiete	19
3.2.4	Supervised Learning	20
3.2.5	K-nearest Neighbor (KNN)	21
4	EMPIRISCHE UNTERSUCHUNG	26
4.1	Implementierung des Artefakts	26
4.2	Angewandte Methodik	28
4.3	Experteninterviews	31
4.3.1	Einleitung	32
4.3.2	Interview-Vereinbarung	32
4.3.3	Fragebogen	32

4.3.4	Vorbereitung und Durchführung der Experteninterviews	33
4.3.5	Nachbearbeitung der Experteninterviews (Transkription)	34
4.3.6	Auswertung der Experteninterviews	34
4.3.7	Hauptkategorien der Inhaltsanalyse	37
4.3.8	Subkategorien der Inhaltsanalyse	39
4.4	Analyse der Experteninterviews	42
4.4.1	Persönliche Informationen	42
4.4.2	Vorwissen im Bereich künstliche Intelligenz	43
4.4.3	Technische Aspekte der KI-Integration	43
4.4.4	Auswirkungen auf den Arbeitsprozess	44
4.4.5	Benutzerfreundlichkeit und Schulung	46
4.4.6	Mitarbeiterrückmeldung und Akzeptanz	48
4.4.7	Herausforderungen und Schwierigkeiten	49
4.4.8	Zukunftsansichten und Entwicklungspotenziale	51
4.4.9	Allgemeine Anmerkungen und Sonstiges	52
4.5	Interpretation der Ergebnisse	53
4.5.1	Persönliche Informationen der Expert*innen	53
4.5.2	Vorwissen im Bereich künstliche Intelligenz	53
4.5.3	Technische Aspekte der KI-Integration	54
4.5.4	Auswirkungen auf den Arbeitsprozess	54
4.5.5	Benutzerfreundlichkeit und Schulung	55
4.5.6	Mitarbeiterrückmeldung und Akzeptanz	55
4.5.7	Herausforderungen und Schwierigkeiten	56
4.5.8	Zukunftsansichten und Entwicklungspotenziale	57
4.5.9	Allgemeine Anmerkungen und Sonstiges	58
4.6	Handlungsempfehlung	58
4.6.1	Schulung und Einführung	58
4.6.2	Anpassungen im Ticketsystem	59
4.6.3	Verbesserung der Benutzerfreundlichkeit	60
4.6.4	Kommunikation und Feedback	60
4.6.5	Zukünftige Entwicklungen	61
4.6.6	Fortlaufende Evaluierungen	61
5	CONCLUSIO	63
5.1	Schlussfolgerung	63
5.2	Ausblick	64

ANHANG A - INTERVIEWLEITFADEN	66
ANHANG B - INTERVIEW-VEREINBARUNG	67
ABKÜRZUNGSVERZEICHNIS.....	68
ABBILDUNGSVERZEICHNIS	69
TABELLENVERZEICHNIS	70
LITERATURVERZEICHNIS	71

1 EINLEITUNG

Am Beginn der folgenden Masterarbeit wird mit dem Kapitel 1 Einleitung begonnen, in dem die Unternehmensstruktur der Styria IT Solutions GmbH & Co KG veranschaulicht wird. Fortgesetzt wird die Masterarbeit mit einer allgemeinen Vorstellung des behandelten Themas inklusive der Zielsetzung. Abgeschlossen wird das Kapitel mit einer Beschreibung der Vorgehensweise, welche den praktischen Teil dieser Arbeit erklären soll.

1.1 Unternehmenspräsentation

Die Styria IT ist eine Tochtergesellschaft des Medienkonzerns Styria Media Group AG, welcher ein Headquarter in Graz und weitere Standorte in Wien und Klagenfurt besitzt. Weiters ist der Konzern auch in Slowenien und Kroatien tätig, was aber in diesem Zusammenhang nicht vertieft wird, da auch die IT-Leistungen in diesen Ländern von einer eigenen IT in Kroatien abgewickelt werden. Abbildung 1 soll das Organigramm der Styria IT veranschaulichen:

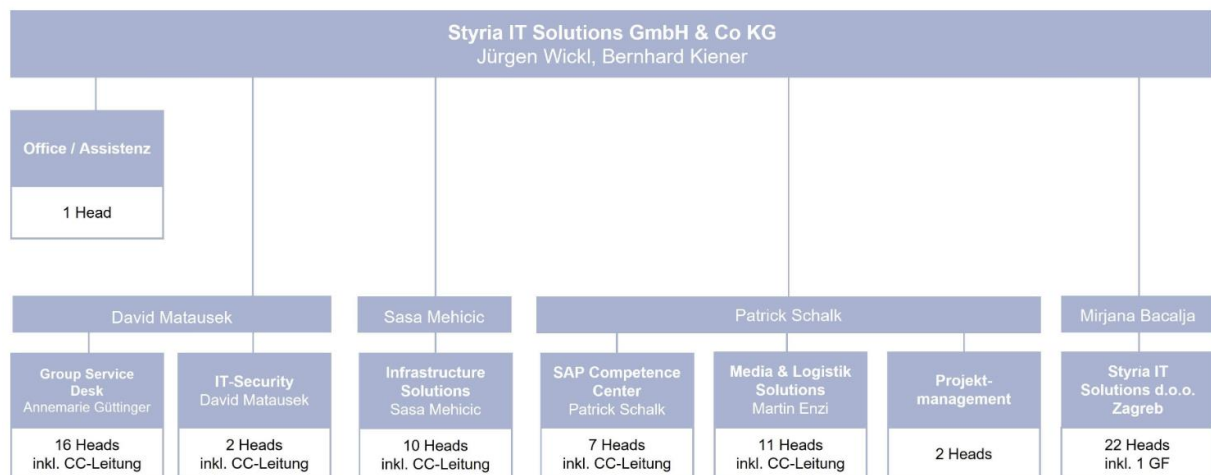


Abbildung 1: Organigramm der Styria IT Solutions GmbH & Co KG (in Anlehnung an (Styria IT, 2023))

In der vorliegenden Arbeit beschäftigen wir uns konkreter mit dem Competence Center „Group Service Desk“, welches wie in Abbildung 1 ersichtlich, als größtes aller Competence Center hinsichtlich der Anzahl der Mitarbeiter*innen in der Styria IT dargestellt wird.

1.2 Ausgangssituation der Masterarbeit

Das Bearbeiten von Serviceanfragen im Kundensupport ist einer von vielen wichtigen Faktoren in der IT. Das Incident Management ist dabei ein zentraler Prozess der IT-Infrastruktur, der die Erfassung, Kategorisierung, Priorisierung, Zuweisung, Eskalation und Lösung von IT-Störungen umfasst. In den letzten zwei Jahren haben Unternehmen vermehrt auf den Einsatz von künstlicher Intelligenz gesetzt und diese nach und nach in Ihre aktuellen Prozesse verankert (Frick et al., 2019).

So auch die Styria IT welche im Service Desk damit startet. Aktuell ist die wirtschaftliche bzw. personelle Lage sehr schwierig und dadurch herrscht auch, sowie in anderen Branchen, in der IT ein Mangel an Fachkräften. Dadurch hat es auch die Styria IT schwer, kompetentes und geschultes Fachpersonal zu finden und auch einstellen zu können. Dieses Problem und auch eine derzeit stattfindende hohe Fluktuation sind Auslöser dafür, dass im Kund*innensupport auf Alternativlösungen zurückgegriffen werden muss (Güttinger, 2023).

Eine resultierende Lösung daraus ist der Einsatz von künstlicher Intelligenz in einigen Prozessen, die ein großes Potential an Automatisierung haben. Dabei soll laut Güttinger (2023) die künstliche Intelligenz helfen, die zeitintensive Prozesse zu erleichtern und den Mitarbeiter*innen mehr Zeit für andere Tätigkeiten ermöglichen. Im aktuell eingesetzten Ticketsystem wird jedes eingegangene Ticket von Mitarbeiter*innen im Service Desk händisch durchgesehen und im Anschluss kategorisiert bzw. dem richtigen Team zugewiesen. Diesen Prozess soll zukünftig eine Lösung auf Basis von künstlicher Intelligenz ablösen und die eingehenden Tickets automatisch klassifizieren und dem jeweiligen Team zuweisen. Im ersten Schritt wird die Lösung folgend implementiert, dass Mitarbeiter*innen im Service Desk auf Basis der vorangegangenen Tickets drei Kategorien vorgeschlagen bekommen, welche mit einer Bewertung versehen sind. Die Mitarbeiter*innen wählen eine der vorgeschlagenen Kategorien aus, oder weisen im schlechtesten Fall die richtige Kategorie aus Erfahrung selbst zu.

In einem weiteren Schritt soll die Prozessautomatisierung erweitert werden, indem zu einem aufgegebenen Ticket, die Datenbank der vorhergegangenen Tickets durchsucht wird und der Lösungsvorschlag von bereits gelösten Tickets, welche vom Inhalt her sehr ähnlich sind, ausgegeben. Im Idealfall wird das Ticket selbstständig von den Kund*innen aufgrund des Lösungsvorschlages gelöst. Ein letzter aktuell möglicher Anwendungsfall ist ein Chatbot, welcher die interne Wissensdatenbank lernen und den Service Desk Mitarbeiter*innen helfen soll, mit gezielten Fragen eine bereits dokumentierte Lösung zu finden. Auch hier kann der Automatisierungsgrad erhöht werden, indem die Kund*innen den Chatbot selbst benützen und dieser aufgrund der Wissensdatenbank einen Lösungsvorschlag ausgibt (Güttinger, 2023). Dieser Prozess wird allerdings im Rahmen dieser Masterarbeit nicht näher betrachtet.

1.3 Zielsetzung der Masterarbeit

Ziel dieser Masterarbeit ist es, dass ein theoretischer Einblick in die Themen Incident Management und künstliche Intelligenz gegeben wird. Dieser Einblick soll die Themen im theoretischen Teil so behandeln, dass ein Grundverständnis für den vorliegenden Praxisteil vorhanden ist und die erzielten Ergebnisse nachverfolgt werden können. Der Praxisteil wird in Kapitel 4 abgearbeitet und analysiert, gefolgt von einer Zusammenfassung und einem Ausblick in Kapitel 5. Ziel ist es weiters, die zu Grunde liegende Forschungsfrage *„Wie beeinflusst die Verwendung von künstlicher Intelligenz die Effizienz und Arbeitsbelastung des Supportteams in einem Ticketsystem?“* zu beantworten und ausreichend zu analysieren. Es wird erhofft, dass die Automatisierung des Prozesses so weit hilft, dass die dadurch gewonnene Zeit der Mitarbeiter*innen für andere relevante Tätigkeiten genützt werden kann.

1.4 Vorgehensweise

Im ersten Schritt soll das Kapitel 2 einen Überblick über die verwendete Methodik geben und im Folgenden soll Kapitel 3 die Theorie wiedergeben und aufzeigen, welche Themenbereiche für die vorliegende Arbeit relevant sind. Der praktische Teil wird ab Kapitel 4 behandelt. Zu Beginn wird die implementierte Lösung, welche vom Implementierungspartner *„Leftshift One Software GmbH“* umgesetzt worden ist, erörtert. Vor der Implementierung der Lösung soll im Rahmen von Experteninterviews ein aktueller Ist-Stand erhoben werden, wie lange die Mitarbeiter*innen für die Bearbeitung bzw. die Kategorisierung der Tickets benötigen. Hier werden verschiedene Service Desk Mitarbeiter*innen an den drei Standorten in Graz, Wien und Klagenfurt befragt. Im nächsten Schritt wird die Lösung im Unternehmen implementiert und einige Zeit von den Service Desk Mitarbeiter*innen verwendet. Mit weiteren Experteninterviews soll ein weiterer Status erhoben werden, wie viel Zeit die Kategorisierung der Tickets mit der Unterstützung der künstlichen Intelligenz in Anspruch nimmt.

Die Befragung der Service Desk Mitarbeiter*innen soll dabei in Form von persönlichen Meetings abgehalten werden. Sollte dies aus irgendwelchen Gründen nicht möglich sein, kann die Befragung alternativ auch über die Kommunikationssoftware Microsoft Teams erfolgen, welche in der Styria IT eingesetzt wird. Als Struktur wird dabei ein Fragenbogen verwendet, welcher in Anhang A zur genaueren Betrachtung angehängt ist. Der Fragebogen soll dabei Schritt für Schritt durchgegangen werden und sowohl klassische geschlossene Fragen als auch offene Fragen für diverses Feedback enthalten. Wünschenswert sind dabei individuelle Kritikpunkte der Expert*innen, die weitere Diskussionen bzw. Verbesserungen des Prozesses einleiten können. Die geschlossenen Fragen sollen dazu dienen, eine grafische Darstellung der Auswertung zu ermöglichen. Der genauere Ablauf der Befragung wird in der empirischen Analyse dargestellt.

Die Analyse der Fragebögen soll im Anschluss an die Befragung erfolgen und unter Zuhilfenahme von Grafiken werden die Ergebnisse zur besseren Veranschaulichung dargestellt. Etwaige auftretende Kritikpunkte der Anwender*innen sollen an den Implementierungspartner weitergegeben werden, um die eingesetzte Lösung kundenspezifisch anpassen zu können und den Prozess zu verbessern. Anschließend werden die erzielten Ergebnisse zusammengefasst und in Kapitel 4 bzw. 5 eine Handlungsempfehlung bzw. Ausblick auf die Situation im Unternehmen gegeben.

2 METHODOLOGIE

Im folgenden Kapitel wird die verwendete Methodologie, welche im Rahmen dieser Masterarbeit behandelt wird, genauer beschrieben. Es handelt sich dabei um das Thema “Design Science Research Ansatz”. Begonnen wird jeweils mit einer Begriffsdefinition, bevor die Thematik selbst genauer erläutert wird.

2.1 Design Science Research Ansatz

Im Kapitel 2.1 wird das Forschungsdesign dieser Masterarbeit, der Design Science Research Ansatz genauer beschrieben. Begonnen wird mit einer Begriffsdefinition und abgeschlossen wird das Kapitel mit einer Prozessbeschreibung in der gestaltungsorientierten Wirtschaftsinformatik. Der gestaltungsorientierte Ansatz ist neben dem verhaltensorientierten Ansatz eine der beiden Forschungsströme in der Wirtschaftsinformatik. Während die Gestaltungsorientierung im deutschsprachigen Raum mehr verwendet wird, ist im englischsprachigen Raum eher die verhaltensorientierte Wirtschaftsinformatik zu Hause. Aufgrund der Komplexität wird der verhaltensorientierte Ansatz im Rahmen dieser Masterarbeit nicht genauer behandelt. Der Design Science Research Ansatz nach Hevner (Peffers et al., 2007) wird in sechs Aktivitäten unterteilt, welche iterativ durchlaufen werden.

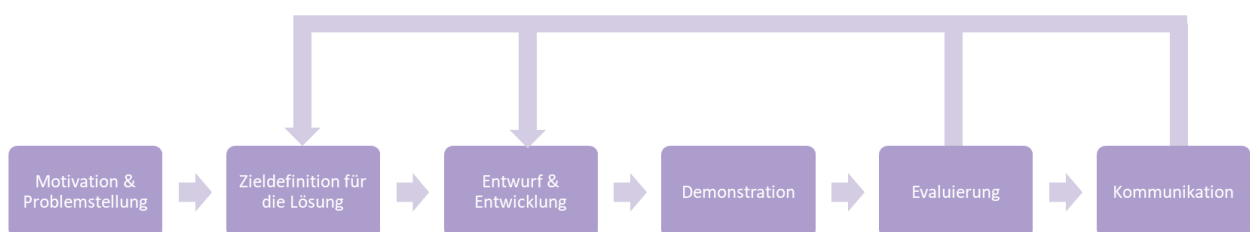


Abbildung 2: DSR-Ansatz nach Hevner (in Anlehnung an (Peffers et al., 2007))

Dieses Forschungsdesign wird im Rahmen der Masterarbeit angewendet und im empirischen Teil genauer beschrieben. Dabei soll auf eine konkrete Problemstellung eingegangen werden, aus dieser eine Forschungsfrage formuliert werden und im Anschluss ein Artefakt entwickelt und eingesetzt werden, welches helfen soll das Problem zu lösen. Der Hauptteil des empirischen Teils dieser Masterarbeit wird die Evaluierung des Artefakts sein, um feststellen zu können, ob das

Artefakt zur Lösung des Problems beigetragen hat. Der Abschluss der vorliegenden Arbeit wird die Darstellung und Kommunikation der Ergebnisse sein.

2.1.1 Begriffsdefinition

Design Science Research (DSR) erklärt die Vorgehensweise, in der ein Artefakt für die Problemlösung erstellt wird. Nach der Implementierung ist es essenziell, dass eine Analysephase stattfindet und die Performance des Artefakts analysiert wird. Das Artefakt ist meist eine technische Lösung, welches das Forschungsproblem beheben soll. Da die Wirtschaftsinformatik sehr anwendungsorientiert ist, ist der Design Science Research Ansatz eine typische Methodologie. Im Idealfall wird das Problem vor dem Einsatz eines Artefakts analysiert und mit den Auswirkungen nach der Einführung des Artefakts verglichen (Marx, 2023).

2.1.2 Gestaltungsorientierte Wirtschaftsinformatik

Die gestaltungsorientierte Wirtschaftsinformatik dominiert den Forschungsansatz im deutschsprachigen Raum. Der Fokus liegt dabei in der Entwicklung und auch der Evaluierung von Lösungen, welche Gestaltungsprobleme in der Wirtschaft und der Verwaltung innovativ, übertragbar und nützlich lösen sollen. In beiden Forschungsansätzen der Wirtschaftsinformatik, sowohl verhaltensorientiert als auch gestaltungsorientiert, spielt die sogenannte Meta-Forschung eine wichtige Rolle. Die Meta-Forschung reflektiert die jeweiligen Prozesse, Ergebnisse und Rollen der Ansätze aus erkenntnistheoretischer und auch aus sozialwissenschaftlicher Sicht. Im gestaltungsorientierten Ansatz wird die Nützlichkeit von IS-basierten Lösungen für wichtige Probleme als akzeptierter Maßstab für Relevanz betrachtet, die Kohärenz bei der Evaluation und Entwicklung kann unterschiedlich sein. Die Kohärenz in der gestaltungsorientierten Wirtschaftsinformatik-Forschung ist aus methodischer Sicht ein wichtiges Thema, jedoch weniger eindeutig definiert und weniger allgemein akzeptiert als die Kohärenz in der verhaltensorientierten Forschung. Methodikfragen, wie ein allgemein akzeptiertes Referenz-Vorgehensmodell für gestaltungsorientierte Wirtschaftsinformatik-Forschung oder die Positionierung und der Charakter von Design-Theorien, sind nach wie vor Gegenstand von Debatten. Es besteht Uneinigkeit darüber, ob "reine" Organisationslösungen ohne IT-Komponenten als geeignete Forschungs- und Entwicklungsgegenstände gelten. Das Fehlen allgemein akzeptierter Evaluationsrichtlinien für verschiedene Artefakttypen zeigt ebenfalls die Notwendigkeit einer stärkeren Fokussierung auf Fragen der Kohärenz in der gestaltungsorientierten Wirtschaftsinformatik-Forschung auf (Winter & Baskerville, 2010).

2.1.3 Prozess in der gestaltungsorientierten Wirtschaftsinformatik

Nach dem Ansatz von Becker (2009), verläuft der Prozess der gestaltungsorientierten Wirtschaftsinformatik in Iteration mit vier Phasen. Diese werden grafisch in Abbildung 3 dargestellt.

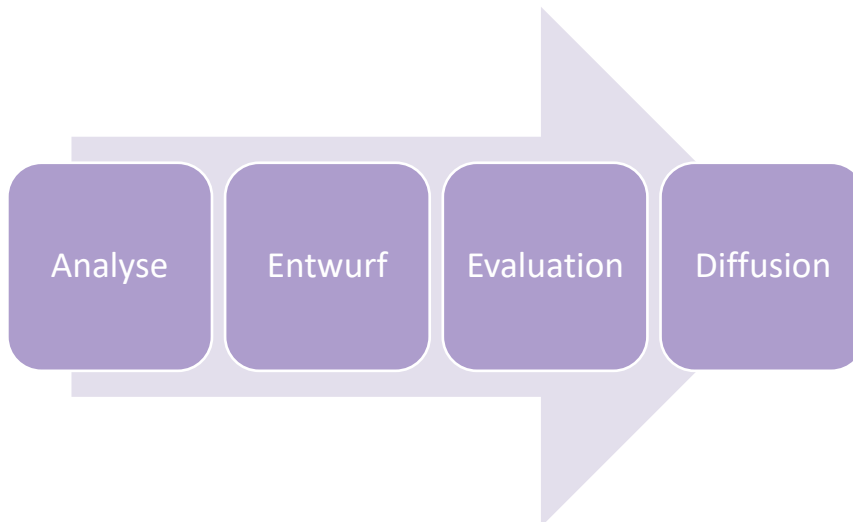


Abbildung 3: Prozess der gestaltungsorientierten Wirtschaftsinformatik (eigene Darstellung, 2023)

Wie in Abbildung 3 ersichtlich startet der Prozess mit einer Analyse des Problems, welches gefolgt durch einen Entwurf eines Artefakts gelöst werden soll. Die Evaluation des Artefakts soll im Anschluss als nächste Phase die Performance des Artefakts aufzeigen, bevor abschließend die Diffusion eingeleitet wird. Die Analyse, als erster Teil des Prozesses, beschäftigt sich mit der Forschung eines Themas aus der Wissenschaft oder anderen Disziplinen. In der Analysephase ist es wichtig, dass eine konkrete Problemstellung definiert wird. Dabei ist es wichtig, dass aus der Problemstellung eine bis mehrere Forschungsfragen bzw. Forschungsziele abgeleitet werden können und die Relevanz der Problemstellung spielt auch eine wesentliche Rolle. Der Blick in die Praxis ist dabei oft hilfreich (Österle et al., 2010).

Wichtig in der Forschung ist die Definition von bestimmten Forschungszielen. Durch die Zielsetzung an sich, lassen sich nach Becker, Holten, Knackstedt und Niehaves (2003) weitere Sub- und Oberziele definieren. Interessant ist dabei die Gliederung in Erkenntnis- bzw. Gestaltungsziele. Erkenntnisziele sind dabei ein Wunsch, gegebene Sachverhalte zu verstehen, um fundierte Vorhersagen über ihre Veränderung zu machen und Gestaltungsziele beinhalten das Verändern von Bestehendem und das Erschaffen von Neuem, unter Nutzung der Ergebnisse aus der erkenntniszielgeleiteten Forschung. Differenziert werden die beiden unterschiedlichen Ziele über den methodischen, bzw. inhaltlich-funktionalen Auftrag. Der methodische Auftrag konzentriert sich auf die Entwicklung von Methoden und Techniken. Diese sollen in Informationssystemen beispielsweise als Beschreibung oder Nutzung dienen. Das Verständnis

und die Gestaltung von Informationssystemen für betriebswirtschaftliche Branchen, wie Handel oder Banken, ist der Hauptaspekt des inhaltlich-funktionalen Auftrags.

Der Forschungsplan spielt in der Analysephase ebenfalls eine große Rolle, da dieser für die Entwicklung und Verbesserung des zu erstellenden Artefakts essenziell ist. Die analysierten Einflussfaktoren des Problems müssen von den Forscher*innen aufgezeigt und in Ihrer Relevanz selektiert werden. Am Ende der Analysephase soll ein geeignetes Forschungsdesign mit einer geeigneten Forschungsmethode entwickelt sein (Österle et al., 2010).

Laut Becker (2009) geht es in der Phase Entwurf darum, dass die Problemstellung aus der Analysephase mit Hilfe der gewählten Methode begründet wird. Die Lösung des Problems soll ein Artefakt darstellen, welches in dieser Phase entwickelt wird. Es muss beachtet werden, dass die Sinnhaftigkeit des Artefakts gegeben ist. Diese kann durch die Grundsätze ordnungsmäßiger Modellierung geprüft werden. In Abbildung 4 werden die Grundsätze aufgelistet und kurz erläutert.

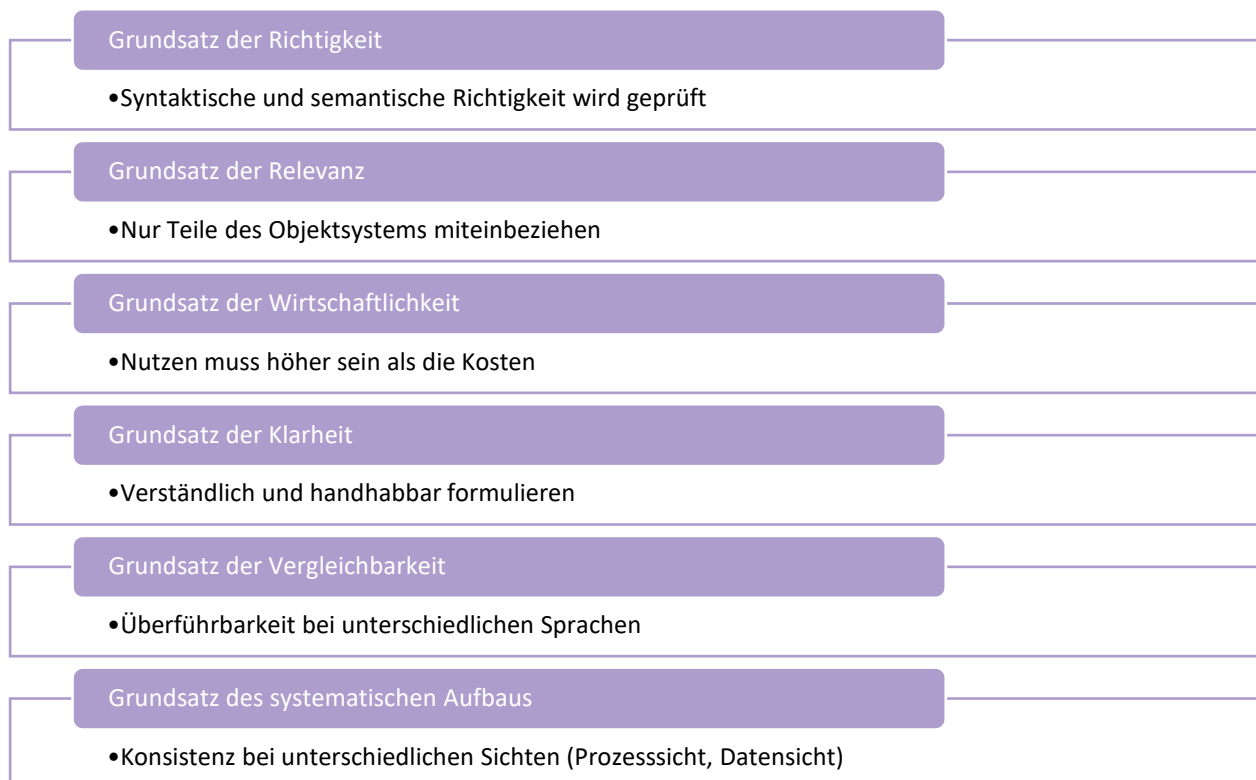


Abbildung 4: Grundsätze ordnungsmäßiger Modellierung (eigene Darstellung, 2023)

Wie in der obigen Abbildung ersichtlich, ist bei der Abstraktion, sprich bei der Detaillierung des Problems bzw. Modells auf diese Punkte zu achten, um ein valides Artefakt entwickeln zu können. Artefakte können beispielsweise Prototypen, Leitfäden, Systeme oder auch Modelle sein, die zur Problemlösung beitragen sollen. Gericke und Robert (2009) beschreiben, dass Artefakte in vier Artefakttypen eingeteilt werden können und sich diese Unterteilung etabliert hat. Die Typen lauten Konstrukt, Modell, Methode und Instanz. Konstrukte beziehen sich dabei grundlegend auf die

Sprache, die Probleme und Lösungen identifiziert und kommuniziert. Modelle repräsentieren Probleme und Lösungsräume, in denen oft Konstrukte eingesetzt werden. Die Methode beschreibt weiters das Vorgehen, wie das Problem gelöst werden kann. Der vierte Typ, nämlich die Instanz versteht die Umsetzung bezogen auf das Problem sowohl von Konstrukten, Modellen als auch Methoden (Becker, 2009).

Die Evaluation, als dritte Phase des Prozesses in der gestaltungsorientierten Wirtschaftsinformatik beschäftigt sich mit der Überprüfung der entwickelten Artefakte. Teile der Überprüfung sind dabei die vorab definierten Ziele und die Methoden des Forschungsplans. Der Nutzen für die Anwenderseite ist dabei eine wichtige Kennzahl, die auf die Allgemeinheit ausgeweitet werden soll. Beispielsweise können im Rahmen der Forschung Kennzahlen wie Effektivität, Effizienz oder auch Nutzen gemessen werden. Oft wird die Evaluation im Rahmen von Reviews durchgeführt. Hierbei muss die Person, welche ein Review durchführt, eine dementsprechende Urteilskraft haben (Österle et al., 2010).

Ein Evaluationsverfahren besteht laut Wilde & Hess (2007) aus einem Objekt, einem Ziel und entsprechenden Evaluationskriterien mit empirischen Metriken oder Indikatoren, welches in Abbildung 5 dargestellt wird.

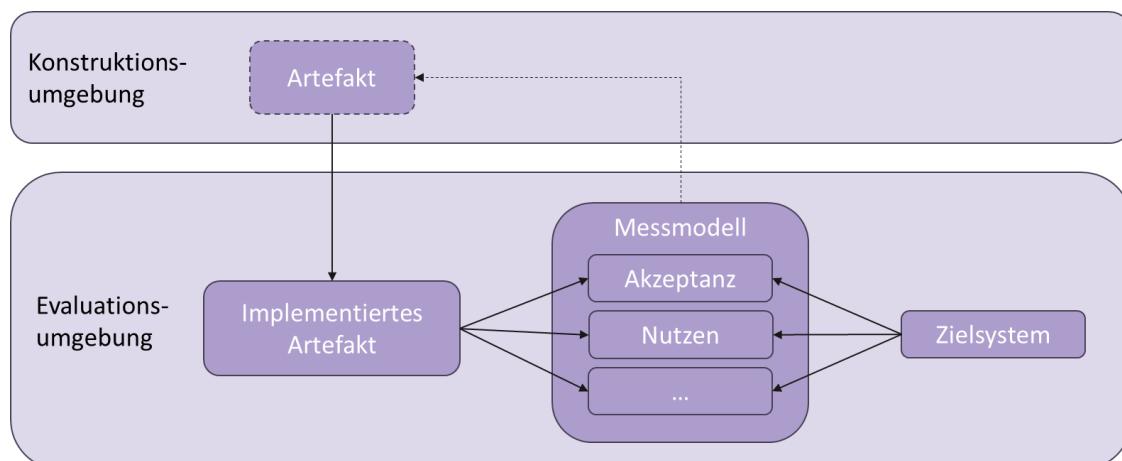


Abbildung 5: Empirische Evaluation von Artefakten (in Anlehnung an Wilde & Hess, 2023)

Die Abbildung soll zeigen, dass das Messmodell in der Evaluationsumgebung einen großen Einfluss hat, da sowohl das Artefakt als auch das Zielsystem Auswirkungen darauf hat.

Die letzte Phase im Prozess ist laut Becker (2009) die Diffusion. Dies bedeutet, dass die Disziplin der Wirtschaftsinformatik bestrebt ist, die Forschungsergebnisse und Erkenntnisse so weit wie möglich unter den Personen und Gruppen zu verbreiten, die Interesse daran haben. Instrumente dafür sind beispielsweise Konferenzen, Dissertationen, Bücher, Vorlesungen, etc.

3 THEORIE

Kapitel 3 beinhaltet die Themen Incident Management und künstliche Intelligenz. Da diese Themen in der vorliegenden Masterarbeit von hoher Relevanz sind, werden diese in den Unterkapiteln definiert und im Detail beschrieben. Das Kapitel künstliche Intelligenz beginnt mit der Geschichte der KI, gefolgt von praktischen Anwendungsgebieten. Im Zusammenhang mit der Umsetzung im Unternehmen, wird in diesem Kapitel weiters auf supervised Learning eingegangen und auch die Methode in der KI K-nearest Neighbor wird nachstehend beschrieben. Die Theorie soll dazu dienen, dass für den praktischen Ansatz ein Grundwissen vorhanden ist.

3.1 Incident Management

Begonnen wird der Theorieteil mit dem Oberbegriff Incident Management. Es wird im Weiteren auf den Zusammenhang mit IT Service Management (ITSM) eingegangen, abgeschlossen wird das Thema mit den Zielen und den Herausforderungen im Incident Management.

3.1.1 Begriffsdefinition

Axelos definiert einen Incident als eine ungeplante Unterbrechung eines Dienstes oder eine Qualitätsverringerung eines Dienstes. Um solchen ungeplanten Ereignissen entgegenwirken zu können, wird in Unternehmen Incident Management angewendet, welches folgend definiert wird. Incident Management Praktiken werden angewendet, um die Zeiträume von ungeplanten Nichtverfügbarkeiten oder Verschlechterungen eines Services zu minimieren. Als wesentlich werden zwei Faktoren genannt, nämlich die frühzeitige Erkennung von Vorfällen und bei unvermeidbaren Incidents muss die schnelle Wiederherstellung des Normalbetriebs gewährleistet sein. Beispiele für einen Incident wären sowohl Ausfälle von Softwarekomponenten oder auch Fehler in der Hardware (Adyrbai, 2021).

3.1.2 Zusammenhang von Incident Management mit ITSM bzw. ITIL

Der Begriff ITSM steht für IT Service Management und umfasst alle Prozesse und Aktivitäten, welche notwendig sind, um die Bereitstellung von IT-Services von Anfang bis Ende gewährleisten zu können. Dabei werden in der IT-Organisation gemeinsame Ziele festgelegt, in welche Richtung die IT-Services gehen sollen. Die Kund*innen sollten dabei miteinbezogen werden, dass die genutzten Services auch für den Kunden zufriedenstellend angeboten werden. Ein weiterer wichtiger Aspekt neben der Bereitstellung der Services, ist der Anwender*innensupport durch den

Service Desk. Dieser muss benutzerorientiert erfolgen, da die individuelle Wahrnehmung der User*innen entscheidend für die Akzeptanz der IT-Services ist. Somit sollen die Quantität und Qualität der IT-Services so geplant, überwacht und gesteuert werden, dass die vereinbarten Ziele und Ergebnisse erreicht werden können. Folgende Kriterien können nach Beims & Ziegenbein (2023) dabei unterstützen:

- Zielorientiert: Die Gestaltung und der Betrieb der Services sollen an den Zielen gemessen werden und sich nach diesen auch orientieren.
- Orientierung nach Geschäftsprozessen: Ein konkreter Beitrag zum Erfolg des Unternehmens ist die bestmögliche Unterstützung der Geschäftsprozesse des Kunden.
- Benutzerfreundlich: Wie bereits erwähnt ist nicht nur die Qualität des Services wichtig, sondern auch die unterschiedliche Wahrnehmung der Anwender*innen. Hochwertige Services allein reichen nicht, sie müssen auch von den Anwender*innen akzeptiert werden.
- Wirtschaftlichkeit: Die Effizienz ist ebenso wie die Effektivität zu betrachten und permanent zu verbessern. Die vereinbarten Ergebnisse sollen somit effektiv erreicht werden, mit angemessenem Aufwand.

Die Verwaltung und Pflegen der IT-Services können in den Unternehmen durch verschiedene Frameworks abgebildet sein. Das wohl bekannteste Framework für die Verwaltung von IT-Services ist die Information Technology Infrastructure Library, auch ITIL genannt. ITIL ist eine herstellerunabhängige Best-Practice-Sammlung, die nicht verbindlich ist. Es basiert auf praktischen Erfahrungen, Modellen und Best-Practice-Architekturen, um professionell IT-Services strukturiert zu betreiben und zu führen (Olbrich, 2008).

Die Implementierung von ITIL bietet zahlreiche Vorteile wie Kosteneinsparungen, verbessertes Risikomanagement und die Optimierung der IT-Betriebsabläufe. Dennoch stehen Unternehmen vor einigen Herausforderungen bei der Umsetzung dieses Frameworks. ITIL-Dokumentationen sind oft unzureichend detailliert und bieten lediglich allgemeine Leitlinien für die Prozesse, die implementiert werden sollen. Viele Manager*innen waren deshalb unsicher darüber, wie sie ITIL am effektivsten einführen sollten, und mussten sich oft stark auf Berater*innen und Dienstleister*innen verlassen. Ein weiteres häufiges Problem bei der ITIL-Implementierung ist der Widerstand der Mitarbeiter*innen aufgrund mangelnder Einbindung in Änderungsprozesse. Um diese Hindernisse zu überwinden oder zumindest zu minimieren, haben Forscher*innen die Nutzer*innenwahrnehmung von IT-Rahmenwerken untersucht. Obwohl ITIL zu einem globalen Standard für Best Practices im IT-Service-Management geworden ist, waren viele Unternehmen der Ansicht, dass die Implementierung von ITIL eine Herausforderung darstellt. Zudem sind nicht alle ITIL-Prozesse für sie von gleicher Relevanz und Wertigkeit (Ahmad & Shamsudin, 2013).

ITIL bietet somit eine Leitlinie für professionelles IT Service Management, und kann unverbindlich angepasst und adaptiert werden. Es beschreibt weiters ein Service Value System, welches nach Randus IT (2022) aus folgenden Dimensionen besteht, um eine Wertschöpfung zu ermöglichen.

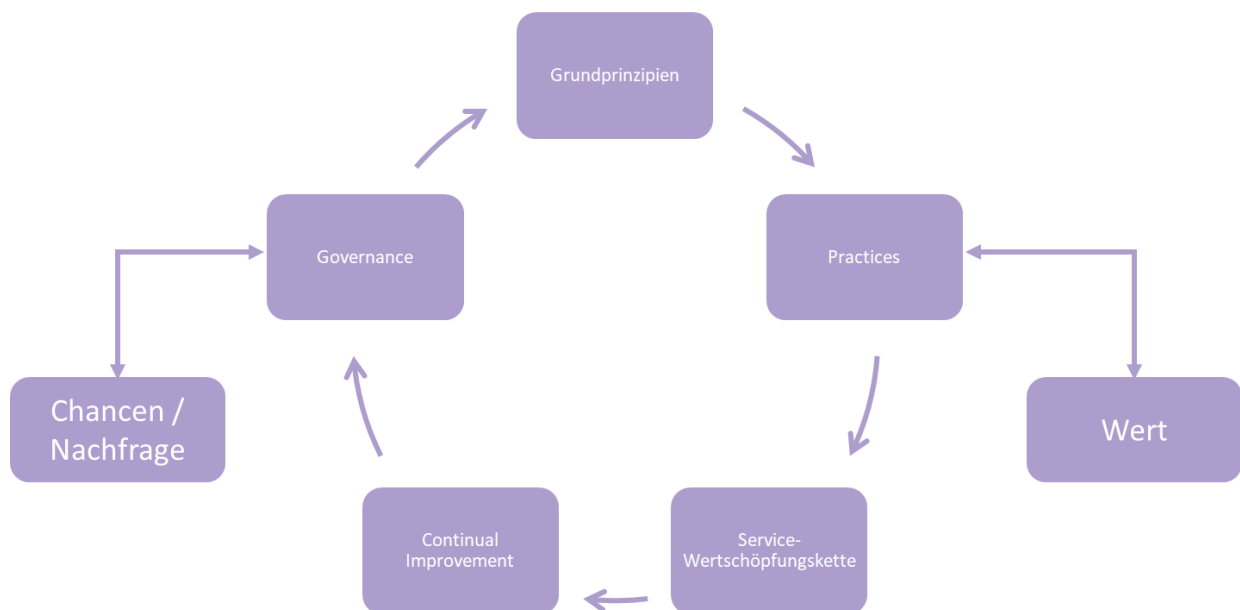


Abbildung 6: ITIL Service Value System (SVS) (in Anlehnung an Randus IT, 2023)

Wie in Abbildung 6 ersichtlich, besteht das Service Value System aus einer Service Wertschöpfungskette, den Management Practices, den Grundprinzipien, der Governance und dem Continual Improvement. Die Grundprinzipien sind dabei die Leitfäden für die Organisation für alle Szenarien. Governance ist die Richtung, wie eine Organisation gesteuert wird. Die Service-Wertschöpfungskette ist eine Reihe an Aktivitäten, um mit einem Service oder einem Produkt einen Mehrwert für Konsument*innen schaffen zu können. Bei den Practices werden Ressourcen zusammengestellt, um gewisse Tätigkeiten abzuwickeln bzw. die definierten Ziele erreichen zu können. Als letzten Punkt ist es wichtig, dass die Wünsche und Erwartungen der Steakholder ständig erfüllt werden. Dies soll die Komponente Continual Improvement sicherstellen. Für jede Komponente des Service Value Systems werden vier Dimensionen betrachtet, welche nachfolgend kurz beschrieben werden (Randus IT, 2022):

- Organisation und Menschen
- Informationen und Technologie
- Partner und Lieferanten
- Wertströme und Prozesse

Die Dimension Organisation und Menschen beschäftigt sich mit den politischen Faktoren, welche beispielsweise die Organisationsstruktur und auch die Kultur beinhalten. Es sollen dabei klar definierte Rollen und Verantwortlichkeiten vorhanden sein, welche unter anderem dafür zuständig

sind, dass eine gewisse Kultur und eine richtige Menge an qualifiziertem Personal vorhanden sind. Durch das Wohlfühlen der Mitarbeiter*innen soll die Serviceorientierung nach Randus IT (2022) sichergestellt werden. Die nächste Dimension Informationen und Technologie beschäftigt sich mit den wirtschaftlichen Faktoren. Diese sollen dazu beitragen, dass die Services mit der dazu benötigten Technologie verwaltet werden können. Technologien sind dabei unter anderem Wissensdatenbanken, Kommunikationssysteme oder auch analytische Tools. Als konkretes Beispiel im Rahmen dieser Masterarbeit ist der Einsatz von künstlicher Intelligenz eine Verbesserung des Service Managements in mehreren Ebenen, wie zum Beispiel in der strategischen Planung des Portfolios hinsichtlich Systemüberwachung bis schlussendlich zum Anwendersupport (Randus IT, 2022).

Partner und Lieferanten ist eine weitere Dimension im SVS. Diese beschäftigen sich konkret mit rechtlichen Faktoren und aufgrund des Zusammenhangs mit der Dimension Organisation und Personen auch mit den Umweltfaktoren. Im Detail bedeutet dies, dass jedes Service bis zu einem bestimmten Grad auch von externen Einflüssen abhängt. Somit beschäftigt sich diese Dimension mit den Beziehungen zu anderen Organisationen, welche an einem bestimmten Service beteiligt sind. Die Kommunikation bzw. das Verhältnis sollte stets positiv gehalten werden, um eine reibungslose Verbesserung als Ziel setzen zu können. Der rechtliche Aspekt wird hier meist durch Verträge oder andere Vereinbarungen festgehalten, so Randus IT (2022).

Wertströme und Prozesse behandelt abschließend die technologischen Faktoren und im Zusammenhang mit den Informationen und der Technologie auch die sozialen Faktoren. Hier werden Steuerungsmechanismen, Prozeduren und auch Workflows definiert, um die vereinbarte Zielsetzung zu erreichen. Es wird dabei kontrolliert, welche Aktivitäten die einzelnen Organisationen durchführt und wie sie sicherstellen, dass die Wertschöpfung effizient und effektiv vorangetrieben wird. ITIL stellt in diesem Kontext ein eigenes Betriebsmodell als Basis zur Verfügung. Wertströme sind dabei aneinandergereihte Schritte, welche die Organisation setzt, um die Services für die Kund*innen zu entwickeln und bereitzustellen. Dabei wird ein klares Bild erzeugt, was die Organisation liefert und wie die Services verbessert werden können (Randus IT, 2022).

Als wichtigster Input für das SVS werden die Chancen und die Nachfrage identifiziert. Die Chancen sind dabei verschiedene Möglichkeiten, um beispielsweise einen Mehrwert für die Steakholder zu schaffen oder auch die Organisation selbst zu verbessern. Die Nachfrage beschreibt direkt den Wunsch oder das Bedürfnis der Kund*innen nach verschiedenen Services oder Produkten. Als Output des Service Value Systems wird ein Wert generiert, der in verschiedensten Formen für die Steakholder*innen generiert wird (Randus IT, 2022).

Andere Frameworks wären beispielsweise „HP IT Service Management“, „IBM Tivoli Unified Process“ oder das „Microsoft Operations Framework (MOF)“. Aufgrund der Vollständigkeit werden

diese Frameworks als Alternativen erwähnt, im Rahmen dieser Masterarbeit aber nicht weiter behandelt (Pröhl & Zarnekow, 2019).

3.1.3 Herausforderungen und Ziele des Incident Managements

Kapitel 3.1.3 beschäftigt sich im Rahmen dieser Masterarbeit über die Herausforderungen und die Ziele des Incident Managements. Wie bereits im vorigen Kapitel erwähnt, bietet das ITIL-Framework eine umfassende Sicht auf alle im Unternehmen eingesetzten Produkte und Services. Die Hauptaufgabe des Incident Managements ist es, dass eingehende Störungen (engl. Incidents) schnellstmöglich behoben werden. Frick et al. (2019) weisen auch darauf hin, dass ein gut geführtes Incident Management Ausfallzeiten der Services reduzieren kann und so die Kund*innenzufriedenheit steigt. Der Prozess der Problemlösung sollte dabei möglichst effektiv und effizient sein, um Störungen, welche gemeldet werden, schnellstmöglich zu beheben. Dabei müssen mehrere Phasen durchlaufen werden, wie beispielsweise die Erfassung der Störung, die Zuweisung zur richtigen Kategorie und schlussendlich die Lösung der Störung. Der Zeitfaktor spielt in diesem Prozess eine essenzielle Rolle und in den meisten Fällen wird der erste Kontakt mit Mitarbeiter*innen aus dem Service Desk bzw. meistens dem First-Level-Support hergestellt. Die weiteren Schritte werden von den Mitarbeiter*innen durch die Kategorisierung abgewickelt und an weitere Strukturen wie Second- oder Third-Level-Support weitergegeben.

ITIL selbst beinhaltet in seinem Framework auch ein entsprechendes Rollenkonzept für die beteiligten Personen. Hierbei wird zwischen folgenden Personen unterschieden:

- Incident Manager
- Mitarbeiter*innen First-Level-Support
- Mitarbeiter*innen Second-Level-Support

Die Hauptaufgabe des Incident Managers ist die Steuerung und Überwachung der Mitarbeiter*innen im Incident Management und der Effizienz. Weiters ist er auch für die Richtlinien des Incident Managements verantwortlich. Der First-Level-Support ist meist die zentrale Anlaufstelle für die eingetretenen Störungen. Zusätzlich zu den oben bereits genannten Tätigkeiten ist auch die Überwachung der Fortschritte der einzelnen Tickets eine Aufgabe des First-Level-Supports. Der Second-Level-Support ist für die weitere Diagnose der Störung verantwortlich und soll mögliche Probleme, die dadurch entstehen können, verhindern. Probleme sind in diesem Fall anders zu definieren als Störungen bzw. Incidents, da Probleme als eine Aneinanderreihung von Störungen im Incident Management definiert werden (Brewster et al., 2016; Frick et al., 2019).

Ein weiterer Aspekt hinsichtlich der Ziele ist die Transparenz der Kosten. Dabei werden messbare Kennzahlen eingeführt, um die Servicekosten beobachten zu können. Mit weiteren Kennzahlen können auch unterschiedliche IT-Organisationen mit dem ITIL-Framework überwacht werden, um deren Erfolg zu messen. Als Beispiel können hier Key Performance Indicators (KPIs) in ITIL für den IT-Support verwendet werden, um die Steuerung der Organisation verwalten zu können (Pröhl & Zarnekow, 2019).

Ein passendes Tool, wie etwa ein Ticket-System soll dabei im Unternehmen eingesetzt werden. Mittlerweile sind Ticketsysteme, welche auch Zugriff auf vergangene Changes, Probleme und bekannte Fehler haben, State-of-the-art und ermöglichen eine schnelle und effiziente Diagnose des Problems. In einem weiteren Schritt wird es in diesen Incident Management Systemen auch möglich, erfasste Störungen automatisch richtig zu kategorisieren und auf bereits passende Analysen zurückzugreifen. Wichtig dabei aus der Unternehmenssicht ist es, dass die Personen, welche im Incident Management tätig sind, immer am aktuellen Stand sein sollen. Dies umfasst nicht nur die Personen im Service Desk, sondern auch die Personen in der Applikationsbetreuung und dem technischen Support. Werden im Unternehmen Produkte bzw. Services von Drittanbieter*innen angeboten, muss mit den Dienstleister*innen eine Vereinbarung bezüglich des Supports abgeschlossen werden. Dabei kann der Support direkt von den Personen des Externen erfolgen oder es wird ein Support des eigenen Service Desk bis zu einem gewissen Grad vereinbart (Randus IT, 2022).

Eine weitere Herausforderung, welche von Zeit zu Zeit immer relevanter geworden ist, sind die Anforderungen von jungen Mitarbeiter*innen. Das Verständnis der IT hinsichtlich Selbstverständlichkeit hat sich in den letzten Jahren grundlegend verändert, da der Umgang mit IT beinahe täglich präsent ist. So werden auch die Anforderungen der Personen immer höher und anspruchsvoller. Die IT-Organisationen müssen sich dabei entsprechend ausrichten, in dem Self Services eingeführt werden, welche einen höheren Wert für die User*innen erzeugen (Pröhl & Zarnekow, 2019).

Der Service Desk leistet aufgrund seiner zentralen Rolle im Incident Management einen großen Beitrag zum Erfolg. Aufgrund der aktuellen Entwicklung hinsichtlich künstlicher Intelligenz, Chat Bots oder verschiedene Visualisierungsmöglichkeiten, können bzw. werden sich die Prozesse im Service Desk in den nächsten Jahren verändern. Dies kann weiters zu einer neuen Ausrichtung dieser Abteilung führen. Die Kommunikation von Störungen veränderte sich ebenfalls in den vergangenen Jahren. Die Vorteile von Social Media werden in dieser Hinsicht verwendet, um über Störungen zu informieren. Jüngere Personen im Service Desk ziehen diese Variante der Kommunikation oft der Kommunikation über E-Mail oder Ticket-System vor (Pröhl & Zarnekow, 2019).

Somit wird der Einsatz von künstlicher Intelligenz auch in diesem Thema erforscht und es wird immer mehr Entwicklungs- und Forschungsarbeit in dieses Thema investiert. Die künstliche Intelligenz soll dabei in bereits bestehende Prozesse integriert und somit optimiert werden. Der Output bzw. die Ziele der Unternehmen bleiben aber meist dieselben. Die Effektivität und die Effizienz, sowie der Mehrwert für Kund*innen soll durch die KI-basierten Services weiter gesteigert werden. Ein weiterer Aspekt, welcher laut Frick, Brünker, Ross und Stieglitz (2019) eingebracht werden kann ist der, dass auch die Fähigkeiten und Leistungen der Mitarbeiter*innen durch die künstliche Intelligenz geschult und verbessert werden können.

Dies ist auch das Ziel der Styria IT Solutions und wird im empirischen Teil dieser Masterarbeit genauer erörtert. Die bereits vorhandenen Prozesse im Incident Management sollen adaptiert werden, in dem durch den Einsatz von künstlicher Intelligenz die Kategorisierung der Tickets automatisiert wird. Die detailliertere Problembeschreibung und die Umsetzung bzw. Evaluierung wird in den nachstehenden Kapiteln erklärt.

3.2 Künstliche Intelligenz

Kapitel 3.2 beschäftigt sich mit der künstlichen Intelligenz (KI) und beginnt mit der Geschichte der KI, gefolgt von praktischen Anwendungsgebieten. Im Zusammenhang mit der Umsetzung im Unternehmen, wird in diesem Kapitel weiters auch auf supervised Learning eingegangen und auch die Methode in der KI K-nearest Neighbor wird nachstehend beschrieben.

3.2.1 Geschichte der künstlichen Intelligenz

Im folgenden Unterkapitel wird kurz auf die Geschichte der künstlichen Intelligenz eingegangen. Viele werden vermuten, dass die künstliche Intelligenz ein sehr neues Forschungsgebiet ist, da diese Materie seit der Einführung von ChatGPT und Ähnlichem einen neuen Hype gefunden hat. Tatsächlich begann die Forschung bereits im Jahr 1950 durch Alan Turing. Sein Ansatz war dabei folgender, dass eine Maschine dann als intelligent gilt, wenn Sie auf blinde Befragungen Antworten liefern kann. Dabei sollte der Anwender nicht unterscheiden können, ob die Antworten von einem Menschen oder einer Maschine getätigt worden sind. Diese Tests wurden später als Turing Tests definiert (Powell, 2019).

Die Turing Tests hatten allerdings nach Logan und Braga (2020) eine essenzielle Schwäche. So kann ein intelligenter Gesprächspartner, der die Befragung durchführt schnell herausfinden, ob die Antworten von einem Menschen oder einer Maschine geliefert werden. Dazu werden persönliche Fragen genutzt, wie beispielsweise „Was ist dein Geschlecht?“, „Waren Sie schon einmal verliebt?“ oder „Welche Sportarten üben Sie aus?“. Somit wird der Turing Test nicht als Intelligenztest, sondern nur als Test eingeordnet, wie gut Programmierer*innen einen Menschen

täuschen können, die Maschine sei auch ein Mensch (Powell, 2019). Der nächste Ansatz in der Forschung wurde 1956 am Dartmouth College in New Hampshire gesetzt. Die Forschung hatte die Vermutung zu Grunde, dass Aspekte des Lernens und die Aspekte der Intelligenz genau so beschrieben werden können, dass es eine Maschine simulieren kann. Die Maschine soll dabei dazu gebracht werden, die Sprache anzuwenden, Konzepte und Abstraktionen zu bilden und im Anschluss die verschiedenen Probleme der Menschen lösen und dabei für zukünftige Abfragen lernen (McCarthy et al., 1955).

Die folgenden Jahrzehnte hatten mehrere Ansätze der Forschung, welche allerdings auf Dauer leider keinen Durchbruch erzielt hatten. Es gab immer wieder viele Stimmen, die gegen den Einsatz von künstlicher Intelligenz gestimmt haben, da der Computer nicht die menschlichen Aufgaben erledigen soll. Weiters gab es auch oft technische oder auch finanzielle Probleme. Der große Durchbruch der künstlichen Intelligenz geschah im 21. Jahrhundert, in dem das Internet eine große Entwicklung machte und riesige Datenmengen in kürzester Zeit zur Verfügung stellte. An diesem Punkt kamen Big-Data-Technologien zum Einsatz und ermöglichten es, diese Daten auch sinnvoll in den verschiedensten Branchen zu nutzen. Die moderne KI ist problemlösungsorientiert und auch datengesteuert. So werden große Datenmengen verarbeitet und beschreiben aufgrund der Vergangenheit, wie bestimmte Probleme gelöst worden sind und erstellen auf Basis dieses Ansatzes ein Modell. Dieses Modell dient dazu, bisher unbekannte neue Probleme ähnlicher Art zu lösen. Weiters ist es sehr beachtlich, welche Fortschritte die künstliche Intelligenz im Bereich der Verarbeitung der menschlichen Sprache gemacht hat. Somit können aktuell zu einem vorgegebenen Thema beliebige Texte kategorisiert und zusammengefasst, oder sogar neue Texte zu diesem Thema verfasst werden. Somit scheint der Durchbruch im Bereich von künstlichen Gehirnen geschafft worden zu sein, allerdings ist dies noch immer mit Vorsicht zu genießen, da die KI immer nur innerhalb der zugrunde liegenden Daten arbeiten kann. Je mehr Daten die künstliche Intelligenz für ihre Modelle zur Verfügung hat, umso mehr kann dem Ergebnis auch vertraut werden. Ein riesiger Vorteil ist es aber, dass in einer beachtlichen Geschwindigkeit die Anfragen in neue Muster angeordnet werden (Know-Center, 2023).

Zusammengefasst kann gesagt werden, dass die künstliche Intelligenz im Alltag der Betriebe nicht mehr wegzudenken ist und viele Entscheidungen, die vorher von Menschen getätigt worden sind, übernimmt. Weiters werden Routineaktivitäten, welche normalerweise periodisch von den einzelnen Personen ausgeübt werden, durch künstliche Intelligenz unterstützt, oder gar komplett abgelöst. Wichtige Begriffe im Bereich künstliche Intelligenz sind zum Beispiel Maschine Learning, Data Science oder auch Big Data. Letzterer wird im nächsten Kapitel genauer erläutert, bevor sich die Theorie mit der eingesetzten Technik in der Styria IT beschäftigt (Fischer, 2020).

3.2.2 Big Data

Das Gabler Wirtschaftslexikon definiert den Begriff Big Data folgend: Big Data bezieht sich auf große Datenmengen aus verschiedenen Quellen wie Internet, Mobilfunk, Gesundheitswesen, Verkehr und intelligenten Geräten, die spezielle Lösungen zur Speicherung, Verarbeitung und Auswertung erfordern. Dadurch können Zwecke wie Rasterfahndung, Trendforschung oder Produktionssteuerung durchgeführt werden. Das Datenvolumen steigt dabei stetig an und bietet neue Möglichkeiten, wie auch Risiken für Benutzer*innen oder auch Organisationen (Gabler Wirtschaftslexikon, 2021).

Die Informationen der Daten können dabei in verschiedenster Form vorhanden sein. Beispielsweise können dies organisierte, unorganisierte, aber auch halborganisierte Daten sein. Die Verarbeitung und Analyse dieser Daten kann oft nur durch komplexe Berechnungen passieren, was für Unternehmen oft ein Hindernis darstellt, da herkömmliche Business-Intelligence-Technologien nicht ausreichend sind. Weiters kann Big Data auch aus verschiedensten Quellen stammen und privat oder öffentlich bzw. auch lokal oder global geteilt sein. Aufzeichnungen, Berichte oder auch Protokolle können in dieser Form als Big Data bezeichnet werden. Oft wird dabei Big Data durch die „5 Vs“ charakterisiert, welches für die fünf Hauptattribute steht:

- Volume
- Variety
- Velocity
- Veracity
- Value

Das Volumen (engl. Volume) repräsentiert dabei die Größe der Daten, welche oft mehrere Terabyte betragen. Variety beschreibt die Datenvielfalt, welche im vorigen Abschnitt kurz beschrieben worden ist. Die Daten können unterschiedliche Arten wie organisiert oder unorganisiert sein und verschiedene Informationen liefern. Organisierte bzw. strukturierte Daten können dabei Informationen sein, welche sich in sozialen Daten oder Tabellenkalkulationen befinden können. Unstrukturierte Daten sind beispielsweise sämtliche Aufnahmen von Video, Ton Bilder oder auch verfasste Texte. Die Maschinen benötigen für Analysen Assoziationen, die bei dieser Art von Daten nicht so einfach hergestellt werden kann. Die Geschwindigkeit (engl. Velocity) beschreibt die Geschwindigkeit der Daten, mit welcher sie erstellt und weiterverarbeitet werden. Veracity definiert die Wahrhaftigkeit der Daten. Genauer gesagt wird damit beschrieben, wie zuverlässig und genau die zu Grunde liegende Datenquelle ist. Die Richtigkeit und Zuverlässigkeit der Daten sind essenziell bei der weiteren Verarbeitung. Das fünfte V beschreibt

den Wert der Daten. Die Auswertung von großen Mengen an Daten erhöht das Wissen und liefert einen hohen Wert für die weiteren Schritte (Abdalla, 2022).

Die Analyse der Daten in solch großen Mengen ist oft eine große Herausforderung. So ist es wichtig, dass wertvolles Wissen aus den Daten entsteht und alle verschiedenen Datentypen in verschiedenstem Umfang betrachtet und analysiert werden. So ist es wichtig, dass die Analyse der Daten detailliert gemacht wird. Beispielsweise können medizinische Daten in verschiedensten Formen, wie Texte, CT-Bilder, Aufzeichnungen von Ärzt*innen und Rezepten bestehen. Dadurch wird eine Analyse komplexer und muss sorgfältig erfolgen. Die Größe bzw. die Datenmenge der Daten ist dabei auch ein wichtiger Faktor. Da diese oft im Terabyte-Bereich liegen erfordert der Forschungsbedarf eine umfangreichere Analyse wie bei herkömmlichen Daten (Ye, Zhang & Zhu, 2022).

Die Klassifizierung ist im Bereich Big Data wichtig und beschreibt eine Methode im Bereich des supervised learnings. Die bekanntesten und verbreitetsten Machine-Learning-Klassifikatoren sind „naive Bayes“, „support vector machine (SVM)“, „k-Nearest Neighbor (k-NN)“, „random forest (RF) und die logistische Regression (LR) (Rashid et al., 2020).

Im Rahmen der vorliegenden Masterarbeit werden nicht alle der oben erwähnten Klassifikatoren genauer erläutert. Die Konzentration gilt der Methode „k-Nearest Neighbor“, welche aufgrund der späteren Implementierung im Ticket-System der Styria IT Solutions, wie der Begriff supervised learning, genauer beschrieben wird.

3.2.3 Praktische Anwendungsgebiete

Künstliche Intelligenz hat in den letzten Jahren eine enorme Entwicklung gemacht und ist in sehr vielen Bereichen nicht mehr wegzudenken. So wurden beispielsweise die Algorithmen für verschiedene Lerntechniken, wie zum Beispiel Deep-Learning, stetig weiterentwickelt und auch die neuronalen Netze haben in Verbindung mit künstlicher Intelligenz Fortschritte zu vermelden. Dabei wird beispielsweise das „Natural Language Processing (NLP)“ immer öfter genutzt, um die Kommunikation zwischen Mensch und Maschine vorantreiben zu können. Output der aktuellen Forschung sind immer mehr eingesetzte Chatbots, die die Kommunikation zwischen Mensch und Maschine verbessern sollen (Rosati, 2023).

Die Chatbots sind nur ein Anwendungsfall der vielen verschiedenen Gebiete. Weiters wird auch in der Robotik bereits künstliche Intelligenz angewendet, um beispielsweise die Fertigung von Autos zu automatisieren oder auch wie am Beispiel von Amazon, wo viele Roboter die Warenlager verwalten. Die Gesichtserkennung ist auch ein weiterer Anwendungsfall, der immer mehr an Bedeutung gewinnt. So verwenden Anwendungen die Gesichtserkennung für die Authentifizierung von Benutzer*innen bzw. wurde in der Corona-Pandemie die

Gesichtserkennung verwendet, um die Sicherheitsabstände zwischen den Menschen messen zu können. Im Sicherheitsbereich wird der Einsatz von künstlicher Intelligenz ebenfalls immer wichtiger. So werden Spam-Filter mit Unterstützung von KI effektiver und auch die Cyber-Sicherheit setzt auf den Trend der KI. So wird unbekannte Software immer besser als harmlos oder auch schädlich eingestuft. Ein Anwendungsgebiet, welches derzeit noch Fortschritte machen muss, ist die Mobilitätsbranche. So wird schon seit Jahren davon gesprochen, dass selbstfahrende Autos bald vermehrt zum Einsatz kommen werden, es aber immer wieder Rückschläge in der Forschung gibt. Weitere Anwendungsgebiete wären beispielsweise der Einkauf, wie auch das Marketing und der Vertrieb (Wuttke, 2023).

3.2.4 Supervised Learning

Am Beginn dieses Kapitels wird versucht, den Begriff „supervised Learning“ zu definieren und im Rahmen dieser Masterarbeit zu erläutern. Auf Deutsch wird der Begriff als überwachtes Lernen bezeichnet und ist ein Ansatz, bei dem mithilfe gekennzeichneteter Trainingsdaten Informationen über Beziehungen zwischen Eingabe und Ausgabe erfasst werden. Überwachtes Lernen ist eine Kategorie des Deep Learning, das gelabelte Daten benötigt, um neuronale Netzwerke zu trainieren. Beim überwachten Lernen wird der Datensatz in Trainings- und Testgruppen unterteilt, typischerweise etwa 75 % für das Training und rund 25 % für das Testen. Die Maschinenlernsysteme lernen aus den Trainingsdaten und entwickeln einen Algorithmus zur Vorhersage des Outputs basierend auf den Inputs. Die Genauigkeit des Algorithmus wird dann mit dem Testdatensatz bewertet. Dieser Ansatz gewährleistet eine ordnungsgemäße Validierung des Algorithmus, da der Testdatensatz keine Daten enthält, die für das Training verwendet wurden. Mit einer großen Menge an gelabelten Daten kann überwachtes Lernen überzeugende Leistungen bei verschiedenen Aufgaben wie Klassifikation und Regression erzielen. Ziel ist es, dass ein System aufgebaut wird, welches den Output für nachfolgend angelegte Inputs vorhersagen kann. Dies führt weiters zu einer so genannten Klassifizierung, wenn die Ausgabe diskret ist oder zu einer Regression, wenn die Aussage kontinuierlich ist. Kontinuierlich würde bedeuten, dass eine Schätzung von numerischen Werten geliefert wird. Modellparameter unterstützen die Darstellung von Informationen, welche auch möglicherweise Schätzungen sein können, wenn die zu Grunde liegenden Trainingsdaten unzureichend sind. Zum „supervised Learning“ gibt es als Gegensatz auch noch das unbeaufsichtigte Lernen, welches als „unsupervised Learning“ definiert ist. In dieser Methode sind nicht zwingend gekennzeichnete Trainingsdaten erforderlich. Ein Mischansatz aus diesen beiden Varianten existiert ebenfalls, welcher „semi-supervised Learning“ genannt wird. Die beiden letztgenannten Ansätze werden aufgrund der Konzentration auf supervised Learning in dieser Arbeit nicht weiter erläutert. Die

folgende Grafik soll das Grundprinzip des Supervised Learning darstellen (Liu & Wu, 2012; Oronowicz et al., 2023; Rehman et al., 2024).

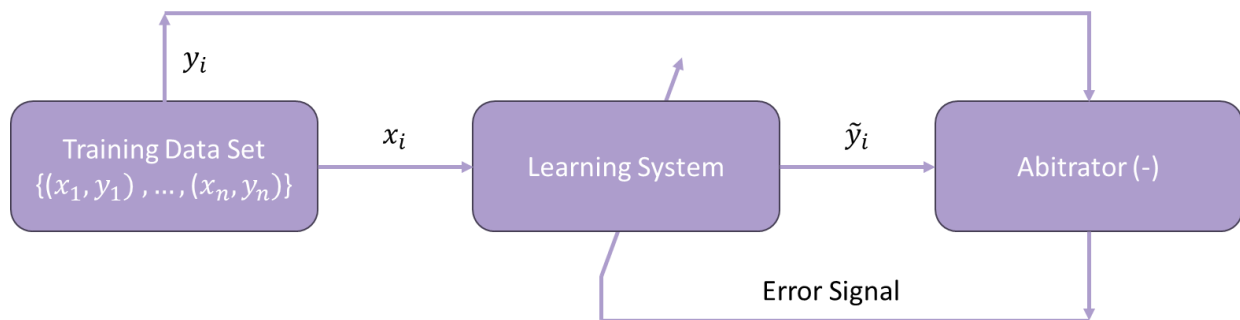


Abbildung 7: Prozess supervised learning (in Anlehnung an Liu & Wu, 2023)

Die obenstehende Grafik zeigt ein Beispielset an Trainingsdaten mit den Parametern (x_i, y_i) , wobei „x“ den Input des Systems repräsentiert und „y“ den Output darstellen soll. Der Output ist dabei die Überwachung bzw. die Beschriftung von Input „x“ und der Index „i“ steht für den Index der Trainingsstichprobe. Im Prozess des supervised learnings wird ein Input „ x_i “ hinzugefügt und das System generiert einen Output „ y_i “. Dieser Output wird dann durch einen Schiedsrichter (engl. Arbitrator) geprüft und mit der Grundwahrheitskennzeichnung „ y_i “ verglichen und die Differenz zwischen ihnen wird berechnet. Die Differenz wird durch den Pfeil „Error Signal“ beschrieben und an das Lernsystem zurückgesendet. Das Lernsystem hat danach die Aufgabe, die Parameter anzupassen und als Ziel soll ein optimales Set an Parametern generiert werden, indem die Differenzen minimiert werden (Liu & Wu, 2012).

3.2.5 K-nearest Neighbor (KNN)

Ein sehr bekannter Algorithmus im Bereich „supervised learning“ ist der Algorithmus K-nearest Neighbor, abgekürzt KNN. KNN ist ein Klassifizierungsalgorithmus im überwachten Lernen, der die Anzahl der benachbarten Datenpunkte (K) verwendet, um Daten zu klassifizieren. Er ist der einfachste maschinelle Lernalgorithmus für überwachtes Lernen und wird häufig für die Lösung von Klassifizierungsproblemen verwendet. Die Klassifizierung dabei dient dazu, um Daten analysieren und Vorhersagen treffen zu können. Diese Methode beinhaltet die Einteilung von Datensätzen in Kategorien oder Klassen anhand verschiedener Parameter. Beispielsweise kann ein E-Mail Provider die E-Mails in verschiedene Gruppen wie Spam, Werbung, Promotion, E-Mails und viele mehr unterteilen (Suyal & Goyal, 2022).

Der KNN-Algorithmus ist eine Methode der Mustererkennung, die Objekte aufgrund ihrer Nähe zu den Trainingsbeispielen im Raum klassifiziert. Sämtliche Berechnungen werden dabei bis zur Klassifizierung aufgeschoben. Der Nutzen dieses Algorithmus kann optimiert werden, wenn wenig bis kein Vorwissen über die Datenverteilung vorhanden ist. Der Algorithmus behält den

ganzen Trainingssatz während des Lernens bei und ordnet den verschiedenen Abfragen eine Klasse zu, die auf einem Mehrheitslabel ihrer k-nächsten Nachbarn basiert. Die einfachste Form des KNN-Algorithmus ist die „Nearest Neighbor rule (NN)“, welche bei $K = 1$ zum Einsatz kommt. Die Proben werden auf Grundlage ihrer Ähnlichkeit mit den benachbarten Proben klassifiziert. Um die Klassifizierung einer bis dato unbekannten Probe vorhersagen zu können, wird der Abstand zwischen der unbekannten Probe und allen Proben im Trainingssatz berechnet. Die Probe mit dem kleinsten Abstand wird als nächster Nachbar betrachtet. Die entstehende Klassifizierung der unbekannten Probe wird auf der Grundlage der Klassifizierung dieses nächsten Nachbarn bestimmt (Imandoust & Bolandraftar, 2013).

Herkömmliche KNN-Methoden zur Textklassifizierung sind sehr rechenintensiv und stark von der Stichprobe abhängig. So muss für die Klassifizierung einer Probe, die Ähnlichkeit mit allen Proben in der Trainingsbibliothek berechnet werden, um den K nächsten Nachbarn zu identifizieren. Passiert diese Textklassifikation mit hochdimensionalen Vektorräumen und einer großen Anzahl von Trainingsmustern, wird die Berechnung die Klassifikationsgeschwindigkeit erheblich beeinträchtigt, was den Nutzen von KNN für die Benutzer*innen schmälert. Die rechnerische Komplexität kann nach Chen (2018) in drei Ansätzen verringert werden:

1. Verringerung der Dimensionalität im Merkmalsraum
2. Minimierung der Berechnung der Musterähnlichkeiten durch Verwendung kleinerer Musterpools
3. Verbesserung der Geschwindigkeit des Algorithmus durch schnellere Suchalgorithmen

Um die Leistung in den verschiedenen Szenarien zu optimieren, wurden verschiedene KNN-Algorithmen entwickelt, welche auf Basis der verschiedenen Datentypen optimiert wurden. Je nach Anforderung und Eigenschaften des Anwendungsfalls, kann beispielsweise zwischen folgenden Algorithmen unterschieden werden:

Name des Algorithmus	Kurzbezeichnung	Unterschiede
Adaptive KNN	A-KNN	• A-KNN ist eine Variante von KNN, die den Wert von „k“ anhand der Charakteristiken der Daten anpasst • Idee dahinter ist, dass die optimale Anzahl von Nachbarn „k“ je nach Daten und Problemstellung variieren kann • Leistung des Algorithmus soll optimiert werden

Locally Adaptive KNN	LA-KNN	• LA-KNN ist eine Erweiterung von KNN, die die Nachbarn lokal anpasst, anstatt eine feste „k“-Anzahl für das gesamte Datenset zu verwenden • Dies bedeutet, dass für jeden Datenpunkt eine unterschiedliche Anzahl von Nachbarn betrachtet wird, basierend auf der Dichte der Daten in der Umgebung
Feature Weighted KNN	F-KNN	• F-KNN ist eine Variante, bei der den Merkmalen eine unterschiedliche Gewichtung zugewiesen wird • Gewichtung wird beeinflusst, je nach dem, wie stark das Merkmal bei der Berechnung der Ähnlichkeit zwischen den Datenpunkten berücksichtigt wird
K-Means KNN	KM-KNN	• Kombiniert Ideen von KNN und K-Means Clustering • Datenset wird in Cluster aufgeteilt, danach wird KNN innerhalb des Clusters angewendet • Effizienz kann in großen Datensätzen gesteigert werden
Weighted KNN	W-KNN	• W-KNN ist eine Erweiterung von KNN, bei der unterschiedliche Gewichtungen für verschiedene Nachbarn berücksichtigt werden • Bestimmte Nachbarn haben aufgrund ihrer Nähe zum Ausgangsdatenpunkt eine höhere Gewichtung als andere Nachbarn

Tabelle 1 : Unterschiede von KNN-Varianten (eigene Darstellung, 2023)

Wie in Tabelle 1 ersichtlich, haben unterschiedliche Varianten von KNN gewisse Vorteile in Ihrer Effizienz und sollten je nach Anwendungsfall angewendet werden (Uddin et al., 2022).

Die Leistung eines KNN-Klassifikators hängt weitgehend von der Wahl von K und der verwendeten Distanzmetrik ab. Die Größe der Nachbarschaft K wirkt sich erheblich auf die Empfindlichkeit der Schätzung aus, da sie den Radius der lokalen Region auf der Grundlage der Entfernung zum K-ten nächsten Nachbarn bestimmt. Kleinere K-Werte können aufgrund von spärlichen Daten und verrauschten oder falsch beschrifteten Punkten zu schlechten Schätzungen führen, während größere K-Werte aufgrund des Einflusses von Ausreißern aus anderen Klassen zu einer Überglättung und verminderten Leistung führen können. Um dieses Problem der

Empfindlichkeit zu lösen, hat sich die Forschung auf die Verbesserung der Klassifizierungsleistung von KNN konzentriert. Die Auswahl des richtigen K-Wertes ist entscheidend und hängt von Faktoren wie der geringen Datenmenge, verrauschten Punkten und dem Vorhandensein von Ausreißern ab. Die Klassifizierungsleistung von KNN ist sehr empfindlich gegenüber dem gewählten K-Wert, da dieser die Schätzung der bedingten Klassenwahrscheinlichkeiten in einer lokalen Datenregion bestimmt. Es wurden gewichtete Abstimmungsmethoden entwickelt, um die Empfindlichkeit gegenüber der K-Wahl zu berücksichtigen, insbesondere wenn die Punkte nicht gleichmäßig verteilt sind. Größere K-Werte sind in der Regel robuster gegenüber Rauschen und erzeugen glattere Klassengrenzen, aber der optimale K-Wert kann bei verschiedenen Anwendungen variieren (Imanodoust & Bolandraftar, 2013).

In Abbildung 8 sind zwei Fälle, welche von Imanodoust und Bolandraftar (2013) beschrieben worden sind, dargestellt: Links mit $K = 1$ und rechts mit $K = 4$. Im Fall von $K = 1$ wird eine unbekannte Probe mit nur einem nächsten Nachbarn klassifiziert, während im Fall von $K = 4$ die vier nächstgelegenen Proben berücksichtigt werden. In beiden Fällen wird die unbekannte Stichprobe der gleichen Klasse zugeordnet wie die Mehrheit ihrer nächsten Nachbarn.

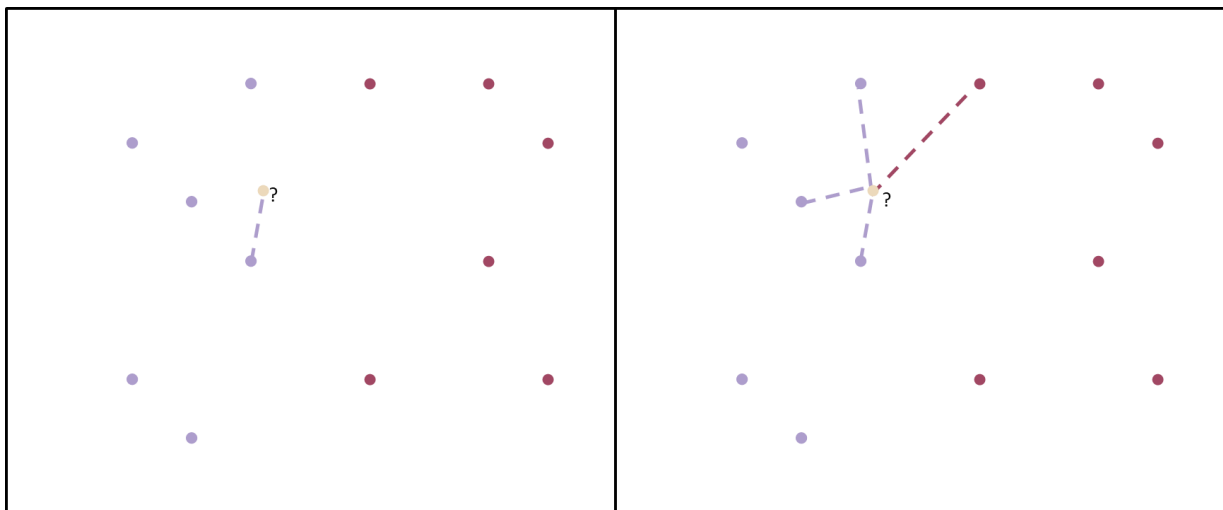


Abbildung 8: Vergleich KNN bei $K=1$ und $K=4$ (in Anlehnung an Imanodoust & Bolandraftar, 2023)

Im Zusammenhang mit der K-Nächste-Nachbarn-Methode (KNN) beginnen wir mit dem 1-Nächste-Nachbarn-Ansatz. Hier suchen wir nach dem Beispiel, das dem Abfragepunkt X (beiger Kreis) am nächsten liegt (violette Kreise). In diesem speziellen Fall ist ein spezifischer Punkt (als Beispiel x_4) das nächstgelegene Beispiel. Das Ergebnis von x_4 (y_4) wird also zur Antwort für das Ergebnis von X (Y) und wir können es wie folgt ausdrücken:

$$Y = y_4$$

Als Nächstes wollen wir die 2-Nächste-Nachbarn-Methode untersuchen. In diesem Szenario suchen wir die beiden nächstgelegenen Punkte zu X, nämlich y_3 und y_4 . Indem wir den Durchschnitt ihrer Ergebnisse nehmen, bestimmen wir die Lösung für Y wie folgt:

$$Y = (y_3 + y_4) / 2$$

Dieses Konzept kann auf eine beliebige Anzahl von nächsten Nachbarn (K) ausgedehnt werden. Zusammenfassend lässt sich sagen, dass bei der KNN-Methode das Ergebnis Y für einen Abfragepunkt X als Durchschnitt der Ergebnisse seiner K nächsten Nachbarn berechnet wird (Imandoust & Bolandraftar, 2013).

Der Output bei der Klassifizierung im maschinellen Lernen wird oft als Label oder Kategorie bezeichnet. Daten können beispielsweise in zwei oder mehrere Kategorien eingeteilt werden, z.B.: die Klassifizierung von E-Mails als Spam oder Nicht-Spam oder die Einteilung von Altersgruppen in Kinder, Jugendliche und alte Leute. Im folgenden Abschnitt werden die einzelnen Schritte des KNN-Algorithmus zur Lösung des Klassifizierungsproblems beschrieben. Essenziell dabei ist, dass der KNN dabei hilft zu bestimmen, zu welcher Kategorie ein neuer Datenpunkt gehört (Suyal & Goyal, 2022).

Folgende Prozessschritte sind nach Suyal & Goyal (2022) im Rahmen des KNN-Algorithmus zu durchlaufen:

1. Die Anzahl „k“ der Nachbarn wird ausgewählt
2. Für die Anzahl „k“ der Nachbarn wird der euklidische Abstand berechnet
3. Gemäß dem berechneten euklidischen Abstand werden die „k“ nächsten Nachbarn ausgewählt
4. Die Anzahl der Datenpunkte jeder Kategorie wird unter den „k“ Nachbarn gezählt
5. Neue Datenpunkte werden den Kategorien zugeordnet, für die die Anzahl der Nachbarn am höchsten ist
6. Das KNN-Modell ist fertig

In der empirischen Untersuchung ist die K-nearest Neighbor Methode Teil der Implementierung. Aufgrund von Umstellungen bei internen Prozessen, kann sich die Klassenzuweisung und die Anzahl der im System verfügbaren Klassen häufig ändern. Weiters kann es oft Mehrdeutigkeiten oder manuell falsch zugewiesene Klassen geben. Für solche Fälle wurde der Algorithmus im System eingeführt. Die Vorgehensweise und eine detaillierte Beschreibung erfolgen in den nachstehenden Kapiteln.

4 EMPIRISCHE UNTERSUCHUNG

Im folgenden Kapitel wird die durchgeführte empirische Untersuchung genauer erläutert. Die Untersuchung dient dazu, die zuvor definierte Forschungsfrage *„Wie beeinflusst die Verwendung von künstlicher Intelligenz die Effizienz und Arbeitsbelastung des Supportteams in einem Ticketsystem“* zu beantworten. Es wird dabei eine qualitative Forschung im Rahmen von Experteninterviews im Rahmen der vorliegenden Masterarbeit durchgeführt. Begonnen wird das Kapitel mit einer genaueren Beschreibung, welche Funktionen das implementierte Artefakt hat und wie es im Unternehmen eingesetzt wird. Bei den Experteninterviews werden die Expert*innen zum Prozess vor und nach der Implementierung befragt, um einen direkten Vergleich herstellen zu können.

4.1 Implementierung des Artefakts

Die eingeführte KI-gesteuerte Ticketkategorisierung stellt einen bedeutenden Fortschritt im Bereich des Ticketsystems dar. Der zugrundeliegende Categorizer, wie er von der Firma Leftshift One betitelt wird, wird über APIs betrieben und zeichnet sich durch seine Fähigkeit aus, eingehende Tickets anhand von drei Vorschlagswerten zu kategorisieren. Ein zentraler Aspekt dieses Systems ist seine Lernfähigkeit, die kontinuierlich durch die Analyse geschlossener Tickets erfolgt. Durch diesen Lernprozess wird der Categorizer befähigt, Muster zu erkennen und seine Kategorisierungsgenauigkeit im Laufe der Zeit zu verbessern (Kloiber, 2023).

Die Implementierung dieses Artefakts lässt sich auch gut in die Perspektiven der Wirtschaftsinformatik nach Myrach einordnen. Es werden verschiedene Perspektiven beschrieben, welche eine Unterschiedliche Interaktion zwischen Mensch und Maschine darstellen sollen. Die Substitutionsperspektive beschreibt dabei das Ersetzen des Menschen durch die Maschine, während die Unterstützungsperspektive die Unterstützung der Maschine für den Menschen beschreibt. In dem Buch werden auch noch andere Perspektiven wie die Kollaborationsperspektive (Mensch interagiert mit Maschine) und die Intermediärperspektive (Spezialist*in wird eingeführt und unterstützt gemeinsam mit der Maschine die Benutzer*innen) (Jung, 2008). Die Einführung von KI in der Ticketkategorisierung entspricht der Unterstützungsperspektive der Wirtschaftsinformatik. Dies bedeutet, dass die KI-Technologie dazu dient, bestehende Prozesse und Aufgaben zu unterstützen und zu verbessern, jedoch nicht notwendigerweise menschliche Arbeitskräfte zu ersetzen. In diesem Fall unterstützt die KI-Technologie Mitarbeiter*innen dabei, Tickets effizienter und genauer zu kategorisieren, anstatt sie vollständig zu ersetzen.

Bei der Bearbeitung neuer Tickets werden dem Benutzer die Top 3 Übereinstimmungen der Ticketkategorien präsentiert, wodurch eine effiziente und nutzerfreundliche Interaktion ermöglicht wird. Beim Öffnen eines Tickets erfolgt eine tiefgreifende Prüfung des Betreffs und der Zusammenfassung durch die KI, die daraufhin die drei höchsten Scores als Vorschlag ausgibt. Diese Scores werden auf einer Skala von 0 bis 1 dezimal berechnet, wobei Wahrscheinlichkeiten eine entscheidende Rolle spielen. Die nachstehende Abbildung zeigt einen Ausschnitt aus dem Ticketsystem, wenn die Service Desk Mitarbeiter*innen ein Ticket manuell aufgrund einer Anforderung bzw. eines Anrufes erstellen (Kloiber, 2023).

← Ticket

Neu

Betreff

Zusammenfassung

Keine Verbindung mit dem Internet möglich

Kategorie

Service Desk

Vorgeschlagene Kategorien

☐	☒ Windows Client	0,975942
☐	☒ Netzwerk	0,579015
☐	☒ VPN Client	0,195506

Abbildung 9: Unklassifiziertes Ticket aus dem Ticketsystem mit Categorizer (eigene Darstellung, 2024)

Wie in Abbildung 9 ersichtlich, werden aufgrund der vorangegangenen Lösungen die drei treffendsten Kategorien aufgrund der Zusammenfassung im obersten Feld vorgeschlagen. Die Mitarbeiter*innen wählen im Anschluss die für sie passende Kategorie aus. Aufgrund der Selektion wird die im Ticket vorgeschlagene Kategorie, welche sich aktuell auf der obersten Ebene mit Namen „Service Desk“ befindet, übernommen und ausgefüllt. Ein Screenshot des klassifizierten Tickets wird in Abbildung 10 dargestellt.

← Ticket

Neu

Betreff

Zusammenfassung

Keine Verbindung mit dem Internet möglich

Kategorie

Service Desk > GSD Servicedesk > Hardware > Windows Client

Vorgeschlagene Kategorien

☒	☒ Windows Client	0,975942
☐	☒ Netzwerk	0,579015
☐	☒ VPN Client	0,195506

Abbildung 10: Klassifiziertes Ticket aus dem Ticketsystem mit Categorizer (eigene Darstellung, 2024)

Durch die konkrete Selektion „Windows Client“ in diesem Beispiel, wird das Feld Kategorie automatisch weiter ausgefüllt und einem Bereich mit Namen „Windows Client“ zugewiesen. Hinter diesem Bereich befinden sich ausgewählte Mitarbeiter*innen, welche aufgrund der Zuweisung eine E-Mail-Benachrichtigung bekommen und sich der Bearbeitung des Tickets widmen können. Der Categorizer schöpft seine Leistungsfähigkeit aus einem supervised Learning Prozess, der durch Natural Language Processing (NLP) unterstützt wird. Der Score, der die Relevanz einer Ticketkategorie widerspiegelt, wird durch die Analyse von natürlicher Sprache ermittelt. Zur weiteren Verfeinerung des Modells wird in der Datenbank nachgelesen, um festzustellen, welche Kategorie in der Vergangenheit korrekt war. Dieses Vortraining stellt sicher, dass der Categorizer die relevanten Kontexte versteht, bevor ein eigener Datensatz für das Feintuning von Texten erfasst wird (Kloiber, 2023).

Laut Kloiber (2023) nutzt der Categorizer für die eigentliche Klassifizierung einen Ansatz ähnlich einem Entscheidungsbaum, bei dem die Klasse vorhergesagt und mit der Wahrscheinlichkeit verknüpft wird. Der Ticketprozess umfasst schließlich die Durchführung durch ein neuronales Netzwerk, welches einen Vektor erzeugt. Dieser Vektor wird mit dem alten Vektor verglichen, um die besten Vorschläge für die Ticketkategorie zu generieren. Insgesamt repräsentiert dieser Ansatz eine ausgefeilte und intelligente Methode zur Automatisierung der Ticketkategorisierung in einem komplexen System.

4.2 Angewandte Methodik

Im Rahmen der vorliegenden Masterarbeit wird ein qualitativer Forschungsansatz in Form von Experteninterviews gewählt. Experteninterviews bieten den Vorteil, dass ein breites Fachwissen aus der beruflichen Praxis abgefragt und im Anschluss repräsentativ für den Anwendungsfall dargestellt werden kann. Die Fragen werden dabei offen gestellt und folgen einem Leitfaden, der vorab definiert wird (Deistler & Rentrop, 2022).

Döring (2023) beschreibt den Ablauf von qualitativen Interviews in zehn Schritten, welche folgend lauten:

Schritt	Bezeichnung	Beschreibung
1	Inhaltliche Vorbereitung	Das Befragungsthema und die Forschungsfrage werden festgelegt. Weiters werden die zu befragenden Personen ausgewählt und ein Leitfaden bzw. der Fragenbogen erstellt. Abgeschlossen ist der Schritt mit einer inhaltlichen Vorbereitung des zu interviewenden Themas.

2	Organisatorische Vorbereitung	Die Vorbereitung der Interviews bzw. der Interviewer*innen ist entscheidend. Wichtig dabei ist der persönliche Kontakt mit den Befragungspersonen. Organisatorisch muss auch diverses Material, wie Audioaufnahmegeräte, Speichermedien, Interviewleitfaden und Informationen über die Forschung gesammelt werden.
3	Gesprächsregeln	Die Gespräche sollen in einer entspannten Atmosphäre stattfinden, indem sich der*die Interviewer*in und die befragte Person zuerst vorstellen und ein wenig mit Smalltalk starten. Danach wird das Interview eingeleitet und auf die Platzierung des Aufnahmegeräts ist dabei zu achten. Eine schriftliche Vereinbarung über Datenschutzmaßnahmen erhöht dabei die Sicherheit bei den Befragten.
4	Durchführung und Aufzeichnung des Interviews	Der*die Interviewer*in hat die Hauptaufgabe, den Gesprächsablauf zu steuern und die eigenen Reaktionen, sowie das nonverbale Verhalten der Befragten zu verfolgen. Interviewer*innen sollen dabei in keinen Situationen zu etwas gedrängt oder verunsichert werden. Bei der offenen Befragung liegt die Verantwortung bei den Interviewer*innen, einen angemessenen Interviewstil zu wählen, wo manche Situationen im Redefluss laufen gelassen werden und auch im Gegenzug in manchen Situationen beim Abschweifen der Person eingegriffen wird.
5	Gesprächsende	Nach dem offiziellen Ende des qualitativen Interviews folgt in der Regel eine informelle Gesprächsphase, die äußerlich der Begrüßungsphase ähneln kann, inhaltlich jedoch meist nicht mit Smalltalk verbunden ist. Während dieser Phase sollten Interviewer*innen trotz möglicher Erschöpfung besonders aufmerksam sein, da Befragungspersonen oft wichtige Informationen nachliefern oder die Gesprächssituation kommentieren.

6	Verabschiedung	Bei der Verabschiedung kann angekündigt werden, dass bei Interesse eine kurze Zusammenfassung der Ergebnisse an die Teilnehmer*innen gesendet wird. Es können beispielsweise während des Interviews Konflikte auftreten, die vorher nicht bewusst waren. Daher ist es auch wichtig, zu beachten, dass ungeplante Interventionen auftreten können. Abschließend soll den Befragungspersonen für die Teilnahme gedankt werden.
7	Gesprächsnotizen	Nach Abschluss des Interviews ist es üblich, Gesprächsnotizen anzufertigen, die oft als Postskriptum bezeichnet werden. Diese Notizen umfassen Beschreibungen der Interviewten, wie deren äußere Erscheinung, seelische Verfassung und Gesundheitszustand. Es können auch eventuelle Unterbrechungen während des Interviews, wie Telefonanrufe oder das Hereinkommen von Kindern dokumentiert werden. Selbst scheinbar Offensichtliches wie das Datum und die Uhrzeit der Befragung sollten notiert werden. Das Postskriptum dient späteren Validitätsbeurteilungen des Materials.
8	Transkription	Die Transkription kann entweder vollständig oder auszugsweise wortwörtlich verschriftlicht werden, muss jedoch vollzogen werden. Heutzutage werden Transkripte aufgrund der digitalen Aufzeichnungen oft mit einer Transkriptions-Software am Computer transkribiert.
9	Analyse der Transkripte	Bei qualitativen Interviews dienen Transkripte der Analyse. Dabei können nicht nur qualitative Analyseverfahren angewendet werden, sondern auch quantitative Verfahren, um eine andere Aussagekraft generieren zu können. Das willkürliche Herausgreifen einzelner Zitate entspricht nicht den wissenschaftlichen Kriterien. Zirkularität ist ein zentrales Prinzip qualitativer Sozialforschung, wodurch Datenerhebung und Datenanalyse in mehreren Zyklen durchlaufen werden.

10	Archivierung des Materials	Durch eine qualitative Befragung entsteht eine große Menge an Material. Diese muss sorgfältig und datenschutzkonform geschützt werden. Interviewäußerungen gelten als persönliche Daten und unterliegen dem Datenschutz. Eine schriftliche Vereinbarung vorab mit dem Befragten kann dabei unterstützen.
----	----------------------------	--

Tabelle 2 - Ablauf von qualitativen Interviews (eigene Darstellung, 2024)

Zur Orientierung wird für die durchzuführenden Interviews ein Leitfaden erstellt. Der*Die Befragte kann auf die offenen Fragen frei antworten und es soll eine möglichst harmonische Gesprächssituation entstehen. Der formulierte Leitfaden dient bei den Interviews dazu, dass das Interview eine gewisse Struktur behält und wichtige Punkte des Interviews berücksichtigt und nicht vergessen werden. Der Inhalt des Leitfadens bzw. der zu stellenden Fragen, kann aufgrund der Gesprächssituation angepasst werden um den Redefluss und das Ziel, ein natürliches Gespräch zu fördern, beizubehalten.

4.3 Experteninterviews

Für die Experteninterviews wurde ein Leitfaden entwickelt, der für eine grobe Struktur der Vorgehensweise dienlich ist. Ziel ist es, dass sechs zuvor definierte Expert*innen des Unternehmens zu dem implementierten Artefakt befragt werden. Um eine natürliche Gesprächsbasis während des Interviews aufrecht halten zu können, soll der Fokus darauf liegen, dass sich die Expert*innen während des ganzen Interviews wohlfühlen und das Gespräch generell, wie ein Gespräch durchgeführt werden soll und kein Prüfungsgespräch oder eine ungute Situation entstehen soll. Wichtig ist es, dass der Prozess der Ticketkategorisierung vor der Implementierung bzw. nach der Implementierung abgefragt wird und ein klares Ergebnis sichtbar ist, wie sich der Prozess verändert hat. Auch ein negatives Ergebnis wird akzeptiert und in Folge analysiert. Das folgende Kapitel soll sowohl die Vorbereitungen inklusive Aufbaus der Interviews und Fragebogen beinhalten wie auch die Durchführung bzw. Nachbearbeitung der Interviews.

Innerhalb der Befragung werden mehrere übergeordnete Themen behandelt, die nachfolgend in Unterkapiteln aufgelistet und kurz beschrieben werden:

- Einleitung
- Interview-Vereinbarung
- Fragebogen

- Persönliche Informationen
 - Technische Informationen
 - Befragung vor der Implementierung des „Categorizers“
 - Befragung nach der Implementierung des „Categorizers“
 - Ausblick
- Abschluss

4.3.1 Einleitung

In der Einleitung werden die Expert*innen begrüßt und es wird erzählt, für was das Interview notwendig ist, bzw. welches Thema behandelt werden soll. Es wird in der Einleitung nochmal explizit erwähnt, dass die Teilnahme der Expert*innen von großer Bedeutung ist und einen wesentlichen Anteil im beschriebenen Forschungsgebiet im Rahmen der Masterarbeit haben. Der zeitliche Aspekt wird auch noch erwähnt und abschließend erfolgt ein Hinweis auf die Datenschutzrichtlinien, welche natürlich im Rahmen der Interviews eingehalten werden. Die Informationen werden vertraulich und anonym behandelt. Um dies auch schriftlich festhalten zu können wird eine schriftliche Vereinbarung, wie in Kapitel 4.3.2 beschrieben, von beiden Parteien unterzeichnet.

4.3.2 Interview-Vereinbarung

Die Interview-Vereinbarung dient dazu, um schriftlich festzuhalten, dass im Rahmen der Experteninterviews datenschutzkonform vorgegangen wird. Gleichzeitig soll es auch für die Expert*innen eine Absicherung sein, dass durch die Befragung keinerlei Nachteile für die Mitarbeiter*innen entstehen. Alle Angaben zur Person oder Institution werden im Nachgang anonymisiert und die Transkription der Interviews dient rein zu Analysezwecken und wird nach Beendigung der Forschung gelöscht.

4.3.3 Fragebogen

Der Fragebogen selbst besteht aus fünf Kategorien, welche eine grobe Struktur des Gesprächs erzeugen sollen. Es muss nicht strikt Frage für Frage abgearbeitet werden, sondern es kann aufgrund der Gesprächssituation eine leichte Abänderung erfolgen. Die persönlichen Fragen dienen dazu, um möglicherweise eine Verbindung zwischen der Altersgruppe und der jeweiligen Erfahrung herstellen zu können. Fortgefahren wird mit den technischen Informationen, wo erfragt

werden soll, welche Erfahrungen mit den beiden Hauptkomponenten der Forschung, nämlich dem Ticketsystem und der künstlichen Intelligenz, bereits vorhanden sind.

Die Befragung zum Abschnitt vor der Implementierung des Artefakts hat unmittelbar vor der Einführung des Artefakts im Juni 2023 stattgefunden. Dort wurde mit den Expert*innen der Prozess eruiert, wie er ohne Hilfe von künstlicher Intelligenz in der Abteilung gelebt wird und wie das persönliche Empfinden dazu ist. Die Befragung nach der Implementierung ist wieder Teil der Experteninterviews, welche im Februar 2024 stattgefunden haben. Dort soll in detaillierterer Art und Weise eruiert werden, wie sich die Lösung im Ticketsystem präsentiert und wie der persönliche Umgang und die Arbeitsweise verändert worden ist. Der abschließende Bereich des Fragebogens ist der Abschnitt Ausblick, welcher in wenigen Fragen einen Blick in die Zukunft geben soll, welche nächsten Schritte im beteiligten Prozess bzw. generell im Ticketsystem gesetzt werden könnten. Das Interview wird mit einem offiziellen Abschluss beendet, indem den Expert*innen erneut für die Teilnahme gedankt werden soll und dass die Ergebnisse der Untersuchung nach der Vollendung für die Expert*innen veröffentlicht werden. Der detaillierte Leitfaden ist in Anhang A sichtbar.

4.3.4 Vorbereitung und Durchführung der Experteninterviews

Im folgenden Kapitel wird sowohl die Vorbereitung als auch die Durchführung der Experteninterviews beschrieben. Es geht einerseits um den organisatorischen Aspekt in der Vorbereitung, als auch um Details zur Ausführung der Interviews mit den Expert*innen.

Im Vorfeld der Experteninterviews wurde der Fragebogen in Zusammenarbeit mit der Teamleitung des Group Service Desks entworfen. Es wurde darauf geachtet, so tief wie möglich ins Detail zu gehen, ohne dabei die technische Komponente der Entwicklung zu beachten. In einem passenden Meetingraum wurde der Fragenbogen bzw. der ganze Leitfaden mit einer Testperson abgearbeitet und eventuelle Fragen bzw. auch Fragestellungen wurden verändert oder sogar aus dem Leitfaden entfernt. Schlussendlich ist der Leitfaden, wie er im Anhang zu finden ist entstanden. Die sechs Expert*innen wurden von der Teamleitung konkret ausgewählt, um eine bestmögliche Mischung aus Expertise zu erlangen. Es wurde eine schriftliche Einladung mit Thema und Vorgehensweise der Befragung an die Teilnehmer verfasst und für alle Interviews der gleiche Raum reserviert. Dadurch sollen technische Probleme mit etwaigem Equipment vermieden werden und es wurde sichergestellt, dass keine anderen Personen die Interviews mithören können. Der Leitfaden wurde bewusst vorab nicht an die Teilnehmer*innen versendet, um vermeiden zu können, dass vorbereitete Antworten gegeben werden und stattdessen emotionale und spontane Reaktionen erleben zu können.

Der Hauptteil der Interviews wurde wie bereits erwähnt im Februar 2024 durchgeführt. Lediglich die Befragung zum Prozess vor der Implementierung des Artefakts fand im Jahr 2023 unmittelbar vor der Implementierung statt. Vier der sechs Interviews fanden direkt vor Ort statt, zwei weitere Interviews wurden über Microsoft Teams abgewickelt. Die Interviews vor Ort wurden mit einem Aufnahmegerät aufgenommen und im Anschluss auch transkribiert. Bei den Interviews mit Microsoft Teams wurde die Aufzeichnung direkt in der Applikation verwendet und im Anschluss ebenfalls transkribiert.

4.3.5 Nachbearbeitung der Experteninterviews (Transkription)

Das Transkribieren ist ein wesentlicher Schritt bei der Auswertung von qualitativen Interviews. Dabei werden die aufgezeichneten Interviews in schriftliche Texte umgewandelt, um sie für die Analyse zugänglich zu machen. Die Genauigkeit ist dabei sehr wichtig, da alle gesprochenen Wörter, Pausen, etc. festgehalten werden müssen. Ein hochwertiges Transkript bildet die Grundlage für eine präzise Analyse, in der verschiedene Muster und Zusammenhänge in den Aussagen der Interviewten identifiziert werden können. Mittlerweile gibt es sehr viele computergestützte Tools – eines der bekanntesten ist die Software MAXQDA. Auch in der vorliegenden Masterarbeit wurde für die Codierung der Transkripte die erwähnte Software herangezogen.

4.3.6 Auswertung der Experteninterviews

Die inhaltlich strukturierende qualitative Inhaltsanalyse nach Udo Kuckartz wird in der Auswertung angewendet. Kuckartz hat eine klare Leitlinie, die er für diesen Ansatz bereitstellt. Seine Methode baut dabei auf der Grundlage von Mayrings Auswertungsmethode auf und wurde von Kuckartz überarbeitet. Somit wurde eine quantitative Darstellung der Analyse mit Unterstützung eines Kategoriensystems ermöglicht. Die strukturierte Inhaltsanalyse hat ihren Fokus typischerweise auf den identifizierten Kategorien und Unterkategorien und deren Beziehungen zueinander. Somit ist diese Methode geeignet für diese Untersuchung. Während der Durchführung der qualitativen Inhaltsanalyse spielen die Kategorien eine zentrale Rolle. Sie werden aus empirischen Daten abgeleitet und dienen als Grundlage für die Analyse. Es ist wichtig anzumerken, dass diese Kategorien eine Vielzahl von Objekten gruppieren und daher eine klare Definition erfordern. Das Kategoriensystem wird zu Beginn aus der Forschungsfrage und der Theorie abgeleitet und anschließend durch Unterkategorien erweitert. Die Inhaltsanalyse besteht aus sieben Phasen, welche nachfolgend beschrieben werden (Kuckartz & Rädiker, 2022):

1. **Initiierende Textarbeit:** Der erste Schritt besteht darin, dass zu analysierende Material sorgfältig zu lesen oder durchzusehen, um ein Verständnis für den Inhalt zu entwickeln.

Dabei werden erste Eindrücke und Ideen gesammelt, die helfen, die Analyse vorzubereiten.

2. **Hauptkategorien entwickeln:** Basierend auf den Forschungsfragen und dem Material, welches zu analysieren ist, werden Hauptkategorien oder Codes entwickelt, die die zentralen Themen oder Konzepte des Textes repräsentieren. Diese Kategorien werden im Voraus festgelegt, um eine strukturierte Analyse zu ermöglichen.
3. **Daten mit Hauptkategorien codieren:** Anschließend wird das Material systematisch nach den vordefinierten Hauptkategorien durchsucht und entsprechend codiert. Passagen im Text, die zu einer bestimmten Hauptkategorie gehören, werden markiert oder mit einem entsprechenden Code versehen. Dieser Schritt wird wie bereits erwähnt unter Zuhilfenahme der Software MAXQDA ausgeführt. Die vorab definierten Kategorien werden im Projekt der Software eingepflegt und den jeweilig zutreffenden Textstellen zugewiesen.
4. **Induktiv Subkategorien bilden:** Neben den Hauptkategorien können während der Analyse auch neue Themen oder Unterkategorien entstehen, die im Material identifiziert werden. Diese werden induktiv gebildet, basierend auf den spezifischen Inhalten und Mustern im Text.
5. **Daten mit Subkategorien codieren:** Nachdem die Subkategorien identifiziert wurden, wird das Material erneut durchsucht, um Passagen zu finden, die diesen Unterkategorien zugeordnet werden können. Auch diese Passagen werden markiert oder codiert, um eine detaillierte Analyse zu ermöglichen.
6. **Einfache und komplexe Analyse:** In diesem Schritt werden einfache statistische Analysen durchgeführt, um die Häufigkeit der Haupt- und Subkategorien im gesamten Material zu ermitteln. Darüber hinaus können auch komplexe qualitative Analysen erfolgen, um die Zusammenhänge und Bedeutungen der identifizierten Themen zu verstehen.
7. **Ergebnisse verschriftlichen:** Schließlich werden die Ergebnisse der Analyse verschriftlicht. Dies kann in Form von Berichten, wissenschaftlichen Artikeln oder anderen Publikationen erfolgen, in denen die wichtigsten Themen, Muster und Erkenntnisse aus der Analyse präsentiert werden.

Zum besseren Verständnis wird die strukturierte qualitative Inhaltsanalyse nach Kuckartz und Rädiker (2022) nachfolgend grafisch dargestellt:

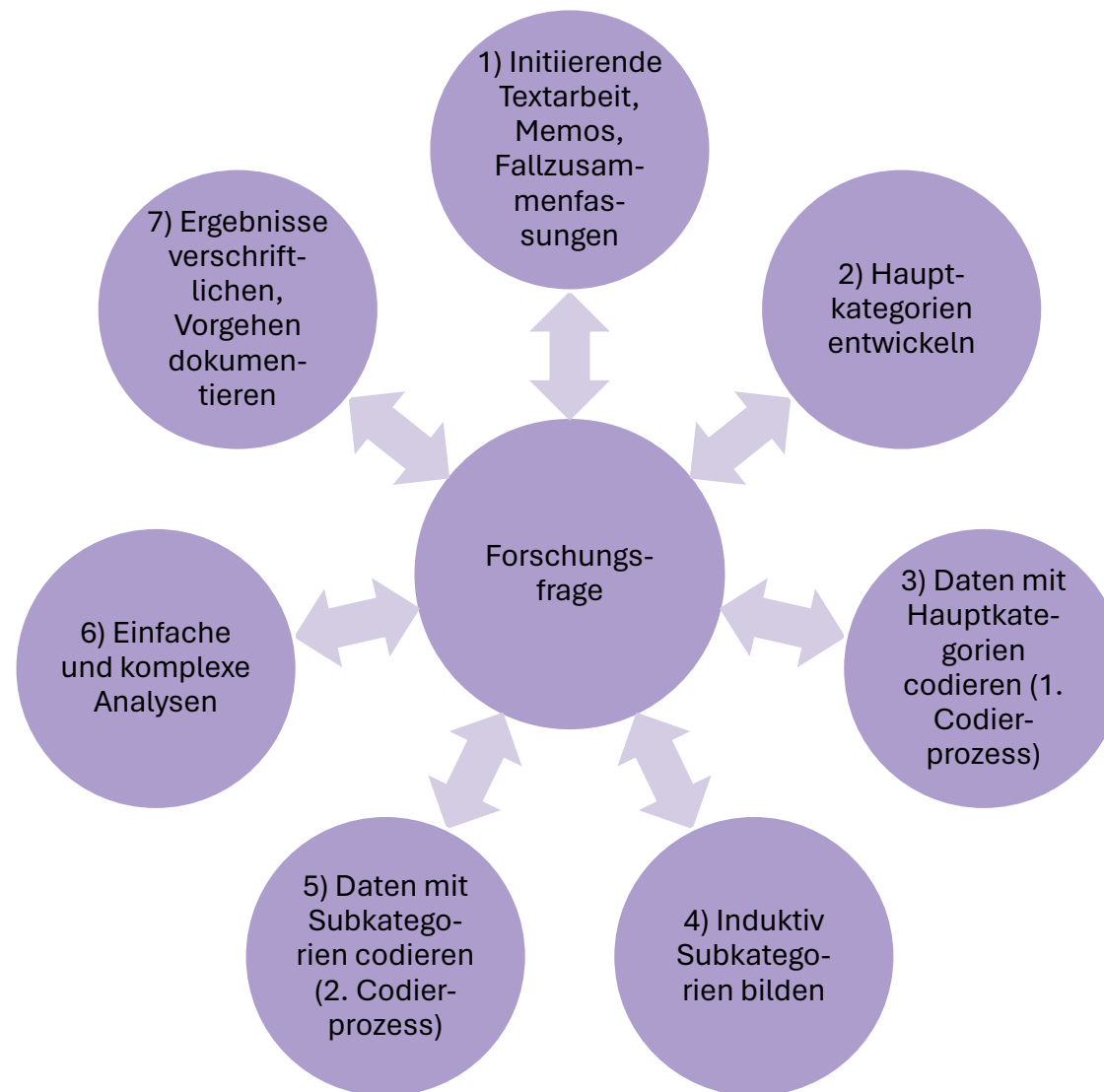


Abbildung 11: Qualitative Inhaltsanalyse nach Kuckartz (in Anlehnung an Kuckartz & Rädiker, 2022)

4.3.7 Hauptkategorien der Inhaltsanalyse

Für die Codierung wurden vorab Kategorien entwickelt, welche hier definiert und genauer beschrieben werden. Wie in den Schritten der Inhaltsanalyse vorab aufgezeigt, wurden Hauptkategorien gebildet, welche die einzelnen Interviews gliedern und strukturieren sollen, um im Nachhinein eine klare Analyse bzw. einen Zusammenhang zwischen den Interviews herstellen zu können. Dabei wurden vor den Interviews, aufgrund der definierten Fragen folgende Kategorien entwickelt:

1. *Persönliche Informationen*

In den persönlichen Informationen soll die Einleitung bzw. der Start der Interviews eingegliedert und sowohl die Altersgruppe als auch die Erfahrung im Bereich Service Desk ausgewertet werden.

2. *Vorwissen im Bereich künstliche Intelligenz*

Da es in der vorliegenden Arbeit auch um den Umgang mit künstlicher Intelligenz geht, wurden die Expert*innen auch über ihr Vorwissen im Bereich künstliche Intelligenz befragt. Dabei soll eruiert werden, inwieweit die Expert*innen sich mit dem Thema beschäftigen bzw. wie sie auch zu der Thematik stehen. Das Vorwissen bezieht sich dabei nicht auf den privaten oder beruflichen Bereich, sondern soll allgemein dargestellt werden.

3. *Technische Aspekte der KI-Integration*

Bei den technischen Aspekten werden zu Beginn alle Informationen der Expert*innen zum Ticketsystem an sich gesammelt. Dabei dürfen auch Kommentare eingeordnet werden, die nicht die Implementierung der künstlichen Intelligenz betreffen, sondern auch die Flexibilität des Systems selbst bei der Anpassung von unterschiedlichen Anforderungen betreffen. Konkret ist die Einführung der künstlichen Intelligenz eine Anforderung an das Ticket-System. Weiters soll hier auch die Genauigkeit bzw. die Zuverlässigkeit der vorgeschlagenen Scores eingeordnet werden. Da die Implementierung von einer externen Firma vorgenommen worden ist und die Expert*innen selbst nicht involviert waren, werden sich die technischen Aspekte und Erfahrungen auf die Erfahrungen in der Nutzung selbst beschränken.

4. *Auswirkungen auf den Arbeitsprozess*

Dadurch, dass sich durch die Implementierung der künstlichen Intelligenz der Prozess verändert hat, sollen hier die Auswirkungen auf den Prozess dokumentiert werden. Hier sollen unter anderem Informationen gesammelt werden, welche den täglichen Arbeitsablauf durch die KI-gesteuerte Ticketkategorisierung betreffen. Weiters wäre es Interessant von den Expert*innen zu erfahren, wie ihre Erfahrungen mit der Interaktion zwischen Mitarbeiter*innen und dem KI-unterstützten System sind. Sollten im Arbeitsprozess durch die Implementierung Herausforderungen oder auch Erleichterungen im Arbeitsprozess durch die vorgeschlagenen Scores auftreten, sollen diese auch in dieser Kategorie erfasst werden.

5. Benutzerfreundlichkeit und Schulung

Ein weiteres wichtiges Thema ist die Benutzerfreundlichkeit des Systems. Das System selbst hat sich im Arbeitsprozess und in der Vorgehensweise verändert, dadurch sollen auch die Einschätzungen dieser Veränderung erfasst werden. Der zweite wesentliche Punkt in dieser Kategorie ist die Erfahrung mit Schulungen oder Anleitungen, die zur Nutzung des KI-gestützten Prozesses unterstützen sollen. Wenn aus dem Gespräch Verbesserungsvorschläge zum Schulungsprozess oder den Anleitungen erwähnt werden, sollen diese auch in dieser Kategorie erfasst werden.

6. Mitarbeiterrückmeldungen und Akzeptanz

In dieser Kategorie soll die grafische Oberfläche außen vorgelassen werden, da diese in einer anderen Kategorie erfasst wird. Konkret geht es in dieser Kategorie darum, dass generell ein Feedback zur KI-gesteuerten Ticketkategorisierung eingeholt werden soll. Die Akzeptanz innerhalb des Service Desks, der betroffenen Abteilung, ist ein weiterer wichtiger Aspekt in dieser definierten Kategorie und soll ebenfalls von den Expert*innen bewertet werden. Es sollen sowohl positive als auch negative Reaktionen auf das KI-unterstützte System identifiziert und niedergeschrieben werden.

7. Herausforderungen und Schwierigkeiten

Die Kategorie Herausforderungen und Schwierigkeiten kann sich in einigen Aussagen mit Kategorie 3 „Technische Aspekte der KI-Integration“ überschneiden, da die technischen Einschätzungen auch oft Herausforderungen oder Schwierigkeiten sein können. Zusätzlich sollen in dieser Kategorie Erfahrungen mit möglichen auftretenden Problemen oder Fehlfunktionen der KI-gesteuerten Kategorisierung erfasst werden. Weiters werden Beschreibungen von Schwierigkeiten oder Herausforderungen bei der Nutzung des Systems in dieser Kategorie erfasst und eventuell, wenn von den Expert*innen erwähnt, auch ein möglicher Lösungsvorschlag des Problems dokumentiert.

8. Zukunftsaussichten und Entwicklungspotenziale

Diese Kategorie wurde für den letzten Teil des Fragebogens erfasst, welcher einen Ausblick über die Zukunft geben soll. Hier sollen einerseits zukünftige Entwicklungen im Bereich der künstlichen Intelligenz im Service Desk betrachtet werden, als auch mögliche Potenziale zur Weiterentwicklung des KI-gestützten Kategorisierungsprozesses eingeordnet werden. Weiters sollen eventuelle Chancen oder Risiken bei der Nutzung fortschrittlicher KI-Technologien identifiziert und erfasst werden.

9. Allgemeine Anmerkungen und Sonstiges

Abschließend wurde eine Kategorie entwickelt, welche offene Kommentare oder zusätzliche Informationen beinhaltet, welche nicht in die zuvor definierten Kategorien passen. Wenn von den Expert*innen in den Interviews erwähnt, sollen überraschende Erfahrungen oder auch unerwartete Erkenntnisse in dieser Kategorie abgebildet werden. Abschließend sollen auch

vertiefende Anregungen für weitere Forschungsrichtungen erwähnt werden. Dies wird vermutlich aber nur der Fall sein, wenn die jeweiligen Expert*innen ein erhöhtes Interesse in der künstlichen Intelligenz besitzen.

4.3.8 Subkategorien der Inhaltsanalyse

Nachdem die Interviews mit der Auswertungssoftware MAXQDA kategorisiert worden sind, ergaben sich für eine feinere Struktur gewisse Unterkategorien, welche folgend beschrieben werden. Für das Verständnis werden auch die Hauptkategorien aus Kapitel 4.3.7 erneut aufgelistet, um schlussendlich einen Gesamtüberblick über die Kategorisierung zu bekommen. Die Subkategorien werden unter den einzelnen Hauptkategorien aufgelistet. Sollten keine Subkategorien notwendig sein, werden alle Wortmeldungen der Hauptkategorie zugewiesen.

1. Persönliche Informationen

- a. Alter: Hier werden die Expert*innen den zur Auswahl stehenden Altersgruppen zugewiesen.
- b. Erfahrung Service Desk: Die generelle Erfahrung im Bereich Service Desk wird hier erfasst.
- c. Erfahrung Styria: Hier wird die Erfahrung der Expert*innen heruntergebrochen auf die Styria IT Solutions kategorisiert.

2. Vorwissen im Bereich künstliche Intelligenz

3. Technische Aspekte der KI-Integration

- a. Ticketsystem allgemein: In dieser Subkategorie sollen alle Meldungen erfasst werden, welche das Ticketsystem selbst und den Umgang damit betreffen. Diese Meldungen müssen keinen Zusammenhang mit der Implementierung der künstlichen Intelligenz zu tun haben.
- b. Technische Aspekte durch KI: Technische funktionelle Aspekte sollen hier mit Zusammenhang der Implementierung erfasst werden. Dies betrifft keine technischen Meldungen hinsichtlich der Implementierung selbst, sondern nur Auffälligkeiten auf Anwender*innenebene.

4. Auswirkungen auf den Arbeitsprozess

- a. Vor KI: Diese Subkategorie soll die erwähnten Auswirkungen vor der Implementierung des Kategorisierungstools aufzeigen. Detaillierter wurde hier auch innerhalb der Subkategorie ein weiterer Code mit Namen „Kategorisierungszeit“ eingeführt, um die Zeit, welche für die Kategorisierung benötigt wird, einschätzen zu können.

- b. Nach KI: Nach KI bündelt die Anmerkungen der Expert*innen zum Arbeitsprozess nach der Implementierung des Tools und wie sich die Arbeitsweise dadurch verändert hat.

5. *Benutzerfreundlichkeit und Schulung*

- a. Benutzerfreundlichkeit: Die Kategorie wird hier in Ihre Bestandteile aufgesplittet und in dieser sollen alle Meldungen zur Benutzerfreundlichkeit erfasst werden.
- b. Schulung: Alle Kommentare hinsichtlich Schulungen werden in dieser Subkategorie erfasst.

6. *Mitarbeiterrückmeldungen und Akzeptanz*

- a. Feedback: Auch diese Kategorie wurde in ihre Einzelteile aufgeteilt. In dieser Kategorie werden alle Rückmeldungen zum Tool selbst gesammelt, welche evtl. nicht die Benutzerfreundlichkeit, sondern eher den moralischen Aspekt der Expert*innen berücksichtigen soll.
- b. Akzeptanz: Diese Subkategorie enthält alle Meldungen, welche explizit auf die Benutzerakzeptanz innerhalb des Teams zurückzuführen sind.

7. *Herausforderungen und Schwierigkeiten*

- a. Vor KI: In dieser Subkategorie werden die Herausforderungen und Schwierigkeiten des Ticketsystems und des Kategorisierungsprozesses vor der Implementierung des Kategorisierungstools erfasst.
- b. Nach KI: Im Vergleich zur vorigen Subkategorie, werden hier die Schwierigkeiten und Herausforderungen eingeordnet, welche es nach der Implementierung aus Sicht der Expert*innen noch gibt.

8. *Zukunftsansichten und Entwicklungspotenziale*

9. *Allgemeine Anmerkungen und Sonstiges*

Wie oben ersichtlich, wurden in den Kategorien 2, 8 und 9 keine Subkategorien gebildet, da diese für nicht notwendig empfunden worden sind, da die Bezeichnung der Hauptkategorie keine spezifische Unterteilung aus Sicht des Erfassers benötigt. Folgend wird eine grafische Übersicht der Kategorisierung veranschaulicht. Ebene 1 zeigt die Hauptkategorien und Ebene 2 würde, wenn definiert, die jeweiligen Subkategorien anzeigen.

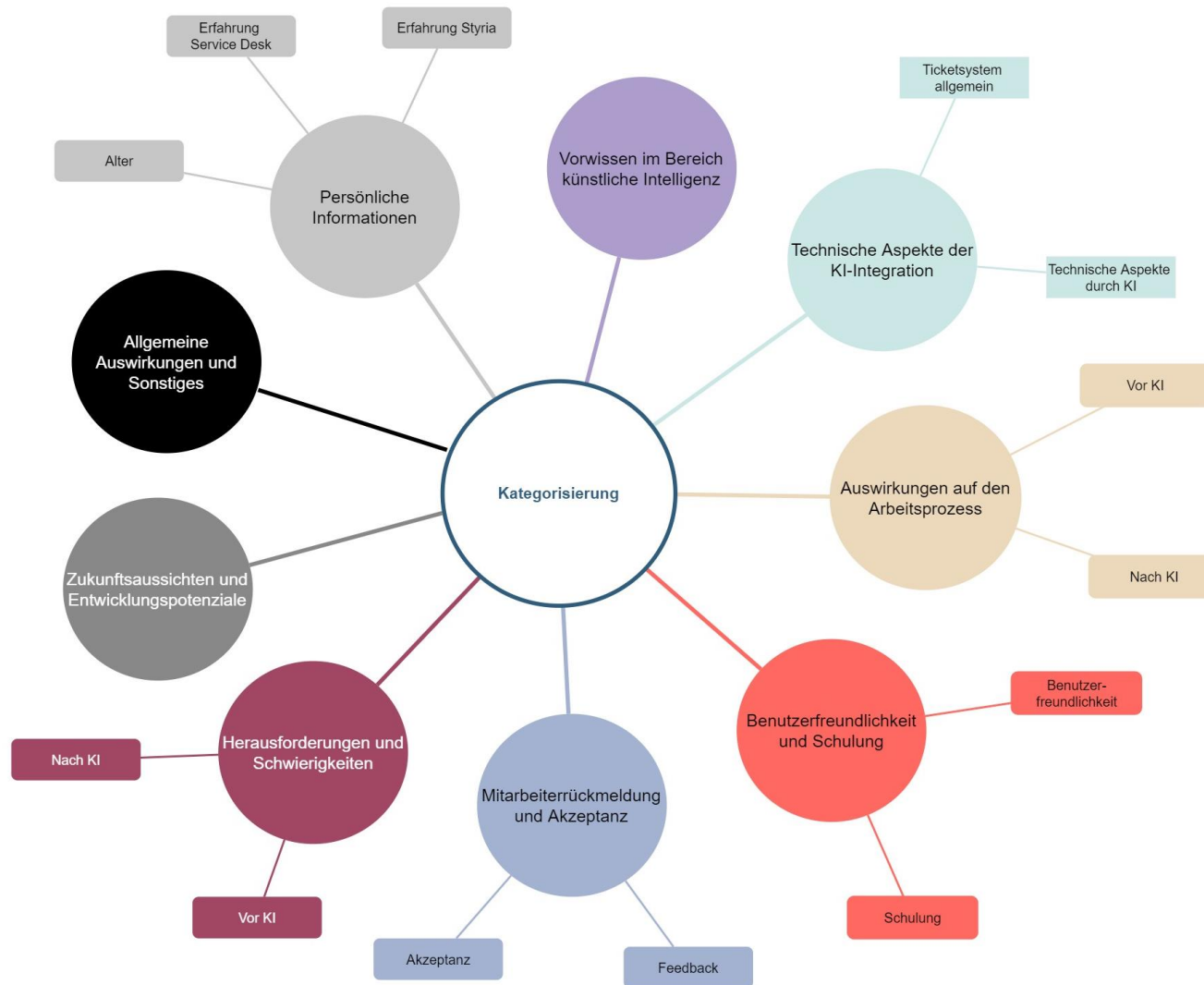


Abbildung 12: Kategorisierung (eigene Darstellung, 2024)

Der nächste Schritt im Vorgehen nach Kuckartz wird die Analyse der mit Codes versehenen Dokumente sein. Ziel dabei ist es, die verschiedenen Kommentare zu einem Code analysieren und Zusammenhänge daraus erzeugen zu können. Die Analyse der Experteninterviews wird im nachfolgenden Kapitel detailliert beschrieben.

4.4 Analyse der Experteninterviews

Bei der strukturierten qualitativen Inhaltsanalyse nach Kuckartz sind die letzten Schritte die Analyse der codierten Dokumente und das Verschriftlichen der Ergebnisse. Diese Phasen werden in diesem bzw. im folgenden Kapitel durchgeführt. Bei der Analyse werden die zuvor verwendeten Kategorien und Subkategorien als Unterkapitel genützt, um eine Gliederung und eine nachvollziehbare Vorgehensweise einhalten zu können.

4.4.1 Persönliche Informationen

Begonnen wird mit den persönlichen Informationen, welche folgend analysiert worden sind. Von den sechs befragten Expert*innen konnten von den zuvor sechs definierten Alterskategorien vier befüllt werden, was für ein breites Spektrum im Service Desk spricht. Lediglich unter 20 Jahren und über 60 Jahren wurden keine Expert*innen befragt. Es gibt sowohl Mitarbeiter*innen die jung und unerfahren sind, aber auch ältere Personen zwischen 50 und 60 Jahren, welche auch einen hohen Wert an Erfahrung mitbringen. Die meisten der interviewten Expert*innen befanden sich im Alter zwischen 30 und 39 Jahren. Nachfolgend wird das Alter im Zusammenhang mit der Erfahrung grafisch dargestellt.

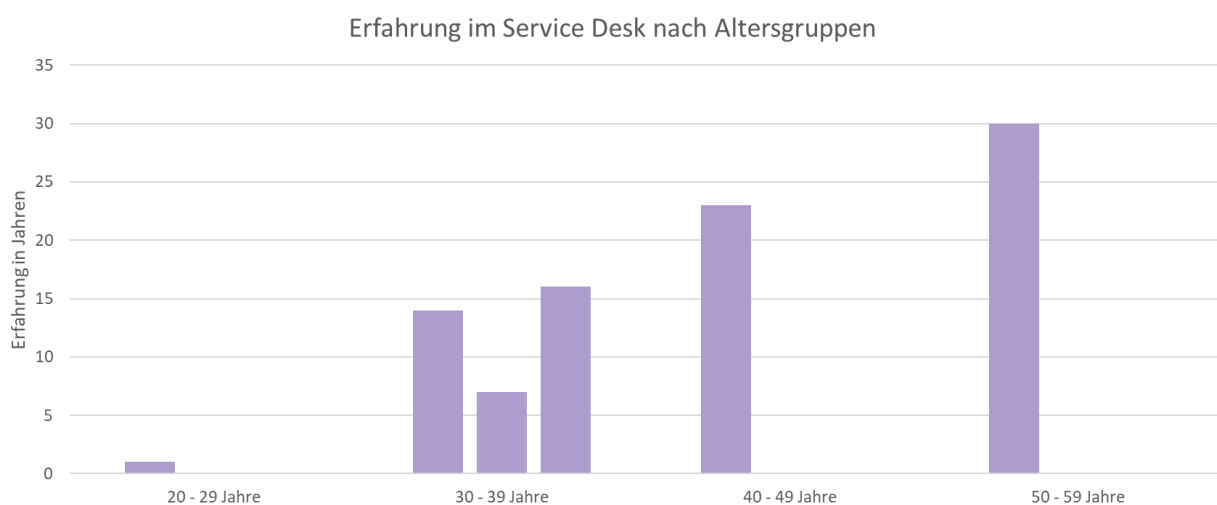


Abbildung 13: Erfahrung im Service Desk nach Altersgruppen (eigene Darstellung, 2024)

Im Detail ist es auch erwähnenswert, dass bei vier der sechs befragten Expert*innen, der Erfahrungswert im Service Desk in Jahren, dem Erfahrungswert im Service Desk der Styria IT

Solutions GmbH entspricht. Lediglich zwei der sechs Mitarbeiter*innen haben davor in einem anderen Unternehmen Erfahrungen im Bereich Service Desk gesammelt.

4.4.2 Vorwissen im Bereich künstliche Intelligenz

Ziel dieser Kategorie war es, das Vorwissen der Expert*innen im Bereich künstliche Intelligenz abzufragen. Vier der sechs befragten Personen gaben bei den Interviews an, dass sie sich wenig bis gar nicht für künstliche Intelligenz interessieren bzw. sich nicht damit beschäftigen. Auch eine versuchte Unterteilung in beruflich bzw. privat führte zu keiner anderen Erkenntnis. Grund dafür ist laut den Expert*innen eine gewisse Skepsis gegenüber der Thematik, die oft auch eine negative Haltung gegenüber diesem Thema einnehmen. Weiters empfinden die Interviewten den aktuellen Stand der künstlichen Intelligenz noch nicht ausgereift genug, um es vermehrt verwenden zu können. Zwei der sechs befragten Expert*innen - im unteren Alterssegment angesiedelt – sagten aus, dass das Thema künstliche Intelligenz sehr spannend sei und auch immer mehr genutzt wird, da es in vielen Bereichen verwendet werden kann. Trotz vermehrter Verwendung fügten die Expert*innen auch hinzu, dass es trotzdem mit Vorsicht zu genießen sei und die derzeitige Entwicklung ein wenig beängstigend ist, da teilweise zwischen Realität und künstlicher Intelligenz nicht mehr unterschieden werden kann. Trotzdem wird es für Standardtätigkeiten wie Bildbearbeitung, etc. im positiven Sinn für gut befunden.

4.4.3 Technische Aspekte der KI-Integration

Im folgenden Kapitel werden die Wortmeldungen zu den technischen Aspekten der KI-Integration analysiert. Wie bereits vorab erwähnt, gibt es hier eine Unterteilung zwischen Kommentaren zum Ticketsystem allgemein und den technischen Aspekten der Implementierung, welche unter anderem die Genauigkeit oder auch die Flexibilität des Systems beinhalten sollen.

Ticketsystem allgemein

In dieser Subkategorie wurden die Kommentare der Expert*innen zum Ticketsystem im Allgemeinen gesammelt. Das Ticketsystem an sich wird durch die Bank für sehr komplex und auch mächtig empfunden. Die Meinungen gehen aber innerhalb der sechs Personen durchaus auseinander. Vereinzelt gefällt den Expert*innen das Ticketsystem wie es ist und hat Potenzial für Entwicklungen, andererseits überwiegen die Wortmeldungen, dass es aktuell eher umständlich ist und es auch im Vergleich zu anderen Ticket-Systemen hintenangestellt wird. Positiv hervorgehoben wird das Arbeiten mit Stichwörtern und auch die vielen Möglichkeiten, welche das Ticketsystem bietet. Bei manchen Expert*innen hält sich die Begeisterung aber auch in Grenzen und es werden Themen wie der Umgang im System und die aktuelle Verwendung

negativ erwähnt. Eine schlankere Version des Ticketsystems wäre bei den Expert*innen wünschenswert und würde vermutlich eher in die Unternehmensstruktur passen.

Technische Aspekte durch KI

Durch die Implementierung der durch die KI unterstützten Kategorisierung haben sich auch neue technische Aspekte entwickelt. Ein wichtiger Punkt im Rahmen der Analyse dieser Subkategorie ist die Erkenntnis, dass die Zeit, welche das Tool für die drei Vorschläge benötigt, für zu lange empfunden wird. Es dauert trotz Eingabe der Pflichtfelder vorab zu lange, bis die KI die drei Vorschläge präsentiert. Weiters ist die Genauigkeit der vorgeschlagenen Kategorien den Expert*innen zufolge auch ausbaufähig. Derzeit gibt es oft Beispiele, bei denen einer der drei Vorschläge fehlerhaft ist und somit ein manuelles Zutun der Mitarbeiter*innen in Form von Kontrolle benötigt. Die Auswahl der vorgeschlagenen Kategorien ist in der aktuellen Umsetzungsform kein Muss, dadurch gibt es auch einige Mitarbeiter*innen, die das Tool nicht verwenden und die Kategorisierung selbst vornehmen. Die vorgeschlagenen Scores waren am Anfang der Implementierung oft abweichend zur tatsächlich richtigen Kategorie. Das Verhalten hat sich aber von Zeit zu Zeit verändert, da die KI von den vergangenen Kategorisierungen lernt, so die Annahme der Expert*innen.

Die Anpassungsfähigkeit im Zusammenhang mit dem Categorizer wird auch als ausbaufähig zusammengefasst, da nach den ersten Vorschlägen von Kategorien, bei einer Änderung im Betreff oder dem Inhalt des Tickets die KI nicht mehr reagiert und ihre davor getätigten Vorschläge nicht anpasst. Ein weiteres Problem ist der zugrunde liegende Kategorisierungsbaum, welcher aktuell verwendet wird. Dieser wird von den Expert*innen als viel zu komplex eingestuft und oft ist es für die Mitarbeiter*innen selbst, auch ohne Unterstützung von künstlicher Intelligenz schwierig, die Tickets richtig zu kategorisieren, da mehrere Kategorien in Frage kommen und hinter den Kategorien teilweise unterschiedliche Abteilungen liegen können. Durch diese Vielzahl an Kategorien ist es oft sehr schwierig eine korrekte Zuteilung zu finden, was oft zu Verwirrung bei den Mitarbeiter*innen führt und dadurch auch zu unterschiedlichen Kategorien eines identen Themas führen kann. Dadurch wird es auch für sehr schwer empfunden, dass die KI hier die richtige Kategorie finden soll, wenn die Mitarbeiter*innen selbst dies nicht zu 100% sicher zuweisen könnten.

4.4.4 Auswirkungen auf den Arbeitsprozess

Die Auswirkungen auf den Arbeitsprozess im Zusammenhang mit dem erneuerten Kategorisierungsprozess durch künstliche Intelligenz, soll in diesem Kapitel beschrieben werden. Zu dieser Kategorie wurden zwei Subkategorien entwickelt, welche die Aussagen der Expert*innen vor bzw. nach der Implementierung darstellen sollen. Es soll hierbei sowohl eine

zeitliche Komponente, sprich wie lange brauchen die Mitarbeiter*innen für die Kategorisierung eines Tickets, dargestellt werden, als auch Themen, die den Prozess aktuell positiv oder negativ beeinflussen.

Vor KI

Der Kategorisierungsprozess, wie er von den Mitarbeiter*innen im Service Desk gelebt werden soll, ist bereits vor der Implementierung der KI nicht 100% definiert und den Expert*innen nicht klar vermittelt worden. Es ging aus den Interviews mehrmals hervor, dass die Mitarbeiter*innen teilweise komplett unterschiedlich arbeiten und in diesem Prozess kein einheitliches Bild vorhanden ist. Auch Verbesserungsvorschläge werden nicht immer aufgenommen, jene die von den Verantwortlichen aufgenommen werden, werden meist auf unbedachte Weise umgesetzt, ohne die Schritte zu testen. Die Umsetzungen sollten dabei mehr durchdacht und auch teilweise Mitarbeiter*innen eingebunden werden. Offiziell gibt es auch eine Vorgabe des Unternehmens, in welcher Zeitspanne ein Ticket kategorisiert bzw. zugewiesen werden muss, worüber nicht alle Expert*innen informiert worden sind und dies somit auch nicht alle berücksichtigen. Solche Vorgaben können auch nur vorgelegt werden, wenn die Mitarbeiter*innen diese auch erfüllen können. Die Expert*innen erwähnten mehrmals, dass oft Tickets auch einfach nur aufgrund der kritischen Zeit kategorisiert und zugewiesen werden, ohne dass sich die Mitarbeiter*innen auch explizit darum kümmern können.

Hinsichtlich der Kategorisierungszeit wurden die Expert*innen befragt, wie lange Sie für die Kategorisierung eines Tickets benötigen. Einig waren sich die Expert*innen über die Tatsache, dass einfache Tickets, in denen sehr schnell hervorgeht um was für ein Problem es sich handelt, relativ schnell kategorisiert werden können. Die Zeitspanne geht hier von fünf bis maximal 20 Sekunden pro Ticket. Komplexere Tickets, welche auch einen langen Mailverlauf zu Grunde haben können, dauern dementsprechend länger. Hier gibt es allerdings auch Probleme, da Mitarbeiter*innen, welche neu im Unternehmen sind, dementsprechend länger brauchen, da diese mit den Strukturen und Verantwortlichkeiten im Konzern noch nicht vertraut sind.

Nach KI

Über die Auswirkungen im Arbeitsprozess wurden die Mitarbeiter*innen auch nach der Implementierung der künstlichen Intelligenz befragt. Die oben erwähnten Auswirkungen vor der KI blieben auch weiterhin bestehen und weiters wurde auch die implementierte Lösung kommentiert. In der aktuellen Form der Implementierung bringt der Prozess keine Veränderung, weder im positiven noch im negativen Sinne. Es muss nach wie vor jedes Ticket von den Mitarbeiter*innen manuell kontrolliert und auch kategorisiert werden. Einfache Tickets, welche auch vorher klar zu kategorisieren waren, sind auch nach wie vor einfach zu kategorisieren und meistens stimmen auch die vorgeschlagenen Scores der KI. Es wurde angemerkt, dass die Vorschläge am Anfang oft fehlerhaft waren und auch teilweise bei gleichen Problemen

auseinander gingen, das hat sich aber mit der Zeit verbessert und kommt aktuell nur mehr vereinzelt vor.

Bei komplexeren Tickets, welche mehrere Probleme in einem Ticket betreffen, oder auch komplexe Probleme, tut sich die KI nach wie vor schwer und bringt keine zufriedenstellenden Scores. Hier wird auch der zugrundeliegende verwendete Kategorisierungsbaum im Service Desk kritisiert, da dieser durchaus der Auslöser für diese Problematik sein könnte. Oft wissen die Expert*innen selbst nicht eindeutig in welche Kategorie das Problem gehören würde, wie soll dies dann eine Lösung mit künstlicher Intelligenz schaffen? Auch bei mehreren Problemen in einem Ticket ist es aktuell so, dass eine Kategorie vorgeschlagen wird, aber es auch andere Kategorien betreffen würde. Hier ist wieder ein großer manueller Aufwand mit der Ticketkoordination verbunden und muss von den Mitarbeiter*innen im Service Desk verwaltet werden.

Die Geschwindigkeit, welche die KI für die Vorschläge benötigt, wird von Expert*innen auch für zu langsam empfunden. Beim Bearbeiten des Tickets dauert es einige Sekunden, bis aufgrund der Informationen die KI die drei Scores vorschlägt. Diese Zeit ist für die Mitarbeiter*innen sehr lange, in dieser Zeit haben Sie das Ticket bereits selbst kategorisiert und benötigen die KI dafür nicht. Für viele hat sich auch aufgrund der Situation, dass diese Vorschläge kein Muss-Feld sind nichts verändert, da trotzdem eigene Kategorisierungen vorgenommen werden können. Somit ist hier ein Tool implementiert, was aber nicht verwendet werden muss und somit auch oft nicht verwendet wird.

Zusammenfassend kann gesagt werden, dass sich der Arbeitsprozess für fünf der sechs befragten Expert*innen nicht verändert hat, da das Tool wenig, bis gar nie verwendet wird und aus verschiedenen Gründen nach wie vor die eigene Einschätzung der künstlichen Intelligenz vorgezogen wird. Trotzdem erwähnen alle Expert*innen, dass die eingesetzte Lösung durchaus Potenzial hätte, wenn man einige Adaptierungen vornehmen würde.

4.4.5 Benutzerfreundlichkeit und Schulung

Das folgende Kapitel handelt über Benutzerfreundlichkeit und Schulung. In dieser Kategorie wurden die Wortmeldungen in die beiden Einzelbereiche als Subkategorien unterteilt und den jeweiligen Textstellen zugewiesen. Die Benutzerfreundlichkeit durch die KI-Unterstützung soll hier analysiert werden, bzw. sollen auch die Erfahrungen mit Schulungen oder Anleitungen eruiert werden.

Benutzerfreundlichkeit

Die Benutzerfreundlichkeit hat sich aus Sicht der Mitarbeiter*innen im Laufe der Zeit verbessert. Es wurden auch bereits kleinere Nachbesserungen im System vollzogen, die es ermöglichen, dass die KI nur aufgrund des Betreffs bereits eine Kategorisierung vornimmt. In der Anfangsphase

war auch eine Beschreibung in einem separaten Feld ein Muss. Somit ist es bei manuell angelegten Tickets bereits eine kleine Verbesserung in die richtige Richtung. Bei automatisch generierten Tickets aufgrund von verschiedenen Kanälen wie E-Mail, etc. dauert der Vorschlag der Kategorien nach wie vor zu lange und somit wird eine manuelle Kategorisierung dem Tool vorgezogen. Die Veränderung im Prozess wird auch von drei Viertel der befragten Expert*innen als nicht erwähnenswert befunden, da es für viele aufgrund der nach wie vor manuellen Kategorisierung nicht relevant ist und die Lösung erst bei schnelleren und genaueren Antworten genützt werden würde. Sollte von den Mitarbeiter*innen eine vorgeschlagene Kategorie zur Vorauswahl gewählt werden, kann diese nicht mehr in Form einer manuellen Kategorisierung verändert werden.

Aus Sicht der grafischen Oberfläche werden die drei Scores in Form eines Ampelsystems mit den Farben Grün, Gelb und Rot gekennzeichnet. Je höher der Score ist, umso mehr geht die Farbe von Rot über Gelb nach Grün. Die farbliche Markierung ist für die Expert*innen sichtbar, was auch so für gut empfunden wird. Teilweise bei einfachen Tickets wird aber wenig darauf geachtet, da die manuelle Kategorisierung schneller durchgeführt wird. Bei komplexeren Tickets, welche zeitlich auch länger für eine Kategorisierung benötigen, wird von der Hälfte der Expert*innen die Kategorisierung als Hilfestellung genutzt. Die andere Hälfte gibt auch der farblichen Komponente und der Vorschläge keine Acht und kategorisiert die Tickets nach wie vor selbst aufgrund der Erfahrung.

Schulung

Das Thema Schulung war bei den Expert*innen ein Thema, welches von allen als kritisch eingestuft worden ist. Die Implementierung und die Einführung des geänderten Prozesses durch die Kategorisierung mit Hilfe von künstlicher Intelligenz wurden ohne Unterstützung der Mitarbeiter*innen durchgeführt. Bei der Einführung wurde auf Anleitung oder ähnliche Materialien verzichtet und das System einfach Live geschaltet. Die Mitarbeiter*innen bekamen im Anschluss die Info, dass die Kategorisierung nun durch das Tool verändert wird und ab sofort genutzt werden sollte. Aufgrund des gerade am Anfang nicht performanten und fehlerhaften Tools, wurde die Einführungsphase als zu kurz empfunden und würde für eine effizientere Implementierung eine genauere Testphase benötigen.

Somit kann gesagt werden, dass die Nutzung des KI-unterstützten Systems in einer gewissen Hinsicht nicht relevant ist, da die Expert*innen das Tool mit dem Hinweis hinnehmen, es sei einfach verfügbar, nutzen tun es aber die Wenigsten aus den zuvor bereits erwähnten Gründen und es nicht genutzt werden muss. Die Schulungsphase fiel laut der Expert*innen auch komplett weg, was einen Umgang mit dem neuen Tool zusätzlich erschwert.

4.4.6 Mitarbeiterrückmeldung und Akzeptanz

Die nächste Kategorie beschäftigt sich mit den Mitarbeiterrückmeldungen und der Akzeptanz. Wie im Kapitel zuvor werden auch hier die beiden Bereiche für die detaillierte Analyse in die Einzelgebiete Feedback und Akzeptanz als Subkategorien unterteilt. Hier sollen in der Subkategorie Rückmeldungen erfasst werden, die idealerweise nicht nur die Benutzerfreundlichkeit betreffen, sondern auch zusätzliche Aspekte hinsichtlich der Implementierung von künstlicher Intelligenz beinhalten.

Feedback

Grundsätzlich kann gesagt werden, dass die Einführung von künstlicher Intelligenz von allen Expert*innen nicht als Verschlechterung angesehen wird. Der Prozess bleibt bei den meisten Mitarbeiter*innen unverändert und wird in der aktuellen Form nur von den Wenigsten als eine Verbesserung angesehen. Hintergrund dazu ist, dass die Lösung bei den Expert*innen nicht als Muss gesehen wird und viele aufgrund ihrer Erfahrung und ihres Wissens, die Kategorisierung selbst vornehmen, da sie in dieser Vorgehensweise meist auch schneller sind. Wenn die Kategorien schneller angezeigt werden würden, finden es die Mitarbeiter*innen grundsätzlich gut, allerdings kann keine Kategorie blind aufgrund ihres Scores verwendet werden. Es muss sehr oft kontrolliert und auch nachgedacht werden, ob der Vorschlag auch der Richtige ist. In vielen Fällen sind die Vorschläge mit hohen Scores auch OK, allerdings gibt es auch immer wieder Vorschläge die nicht zum aktuellen Problem passen. Für neue Mitarbeiter*innen wird die Implementierung für positiv empfunden, da die Vorschläge, auch wenn nicht immer richtig, eine Richtung geben können und somit unterstützen kann.

Akzeptanz

Da der Einführungs- bzw. auch der Schulungsprozess bei den Mitarbeiter*innen nicht präsent war, ist auch die Akzeptanz für die Kategorisierung im ersten Schritt nicht positiv. Wie bereits erwähnt wird es nicht als Nachteil gesehen, trotzdem waren am Anfang bei der Einführung die Expert*innen dagegen und negativ gestimmt. Der Einsatz des Tools hat aber für die Expert*innen keine Auswirkung auf den Arbeitsprozess, da wie bereits erwähnt die Verwendung nicht verpflichtend ist. Somit stößt das Unternehmen hier auch auf keine großen Probleme, da laut den Expert*innen das Tool vorhanden aber Großteils nicht genutzt wird. Die negative Stimmung hat sich laut den Mitarbeiter*innen im Laufe der Zeit auch ein wenig verbessert und die zuerst gegen die Implementierung waren, sind mittlerweile durchwegs neutral gestimmt und versuchen die Kategorisierung ab und zu. Trotzdem gab es recht kritische Äußerungen der Expert*innen, wie etwa, dass diese Implementierung in der aktuellen Form für unnötig empfunden wird und es keine Verbesserung im Prozess gibt. Das Geld und die Zeit hätte den Expert*innen zufolge auch anders

investiert werden können, da es auch andere Problematiken im Service Desk gibt, die früher behandelt werden sollten.

4.4.7 Herausforderungen und Schwierigkeiten

Das Kapitel Herausforderungen und Schwierigkeiten soll die Erfahrungen mit möglichen Problemen oder auch Fehlfunktionen der KI-gesteuerten Kategorisierung aufzeigen. Weiters sollen auch erwähnte Schwierigkeiten und Herausforderungen vor bzw. auch nach der Implementierung beschrieben werden. Mögliche Lösungsansätze für identifizierte Probleme und Herausforderungen sollen in Kapitel 4.5 im Anschluss aufgezeigt werden. In dieser Kategorie gab es wie auch schon in manchen Kategorien davor eine Unterteilung in Rückmeldungen vor Einsatz der KI und Wortmeldungen nach dem Einsatz der KI.

Vor KI

Eine erste große Herausforderung im Kategorisierungsprozess ist die richtige Kategorisierung und die dadurch entstehende Zuweisung zu einem gewissen Team, welche das Problem bearbeiten soll. Eine falsche Kategorisierung hat einen zusätzlichen Aufwand zum Thema, da die Abteilungen das Ticket prüfen müssen und eventuell nicht dafür zuständig sind. Somit liegt es bei den zuständigen Personen im Service Desk, das Ticket genau zu prüfen und überlegen, ob das Problem selbst gelöst werden kann, oder gewisse Kolleg*innen zugezogen werden müssen.

Weiters ist es für neue Mitarbeiter*innen schwer, die Vielzahl an Kategorien richtig zu verwenden und auch zu wissen, wer ist im Endeffekt für die Bearbeitung des vorliegenden Problems zuständig. Oft sind mehrere Kategorien auf den ersten Blick für ein Ticket richtig, hinter den Kategorien können aber verschiedene Abteilungen stehen. Beispielsweise können gewisse Zuordnungen bei User*innen von den Mitarbeiter*innen im Service Desk selbst vorgenommen werden, komplexere Anfragen werden von der Infrastruktur-Abteilung abgewickelt. Die Kategorien im Kategorienbaum zeigen den Unterschied aber oft nicht auf. Ein weiteres Beispiel wäre auch im SAP-Bereich, wo die Zuständigkeit zwischen der Installation der grafischen Oberfläche, oder auch Problemen bei der Anwendung im System in unterschiedlichen Teams liegen und somit unterschiedliche Kategorien ausgewählt werden müssen. Das kann bei Tickets aus E-Mails problematisch werden, wo die Kund*innen möglicherweise keine konkreten Details preisgeben.

Ein weiteres Problem bei den Expert*innen ist die Herausforderung, dass Kund*innen oft mehrere Probleme in einem Ticket melden können. Dies kann automatisch per Mail passieren, da hier aus einer Mail an eine bestimmte Adresse ein Ticket erstellt wird. Somit ist es für die Expert*innen auch schwierig, die richtige Kategorie auszuwählen und auch dafür zu sorgen, dass die anderen Probleme ebenfalls von den zuständigen Personen abgewickelt werden. Auch hier ist laut den

Expert*innen keine einheitliche Vorgehensweise bekannt, wie in diesem Fall damit umzugehen ist. Sehr komplex wird es dann, wenn bereits davor ein langer Mailverlauf zwischen verschiedenen Parteien stattgefunden hat und zu einem späteren Zeitpunkt das Ticket durch eine Mail an den Service Desk erstellt wird. So müssen die Mitarbeiter*innen den ganzen Mailverlauf analysieren und eine passende Kategorisierung treffen, was sehr viel Zeit kosten kann.

Eine weitere Variante für die Ticketerstellung ist das Ticketportal der IT selbst. Hier können Kund*innen über das Portal ein Ticket erstellen, oder in der Wissensdatenbank nach Problemen suchen. Da die grafische Aufbereitung und das Handling nicht sehr ansprechend sind, wird diese Variante bei den Kund*innen sehr wenig bis gar nicht genutzt. Für den Service Desk wäre diese Variante eine Hilfe, da hier das Ticket über ein vorgefertigtes Formular ankommen würde.

Allgemein kann gesagt werden, dass die Expert*innen eine individuelle Vorgehensweise bei gleichen Prozessen kritisieren und hier eine Verbesserung geschehen muss, die eine einheitliche Arbeitsweise voraussetzt. Zusätzlich kommt dazu, dass der Service Desk in der Styria IT Solutions aufgrund von Mangel an Mitarbeiter*innen kein gewöhnlicher Service Desk mit First-Level-Support, Second-Level-Support, etc. ist, sondern die Mitarbeiter*innen oft zwischen diesen Rollen agieren. Somit wird das Kategorisieren auch teilweise bei den Expert*innen als zeitfressend kommentiert und es liegt für die Kategorisierung nicht viel verfügbare Zeit vor.

Nach KI

Ziel ist es durch die Implementierung, dass einige der oben genannten Probleme mit Hilfe der künstlichen Intelligenz behoben werden sollten. Die Herausforderungen und Schwierigkeiten, welche oben bereits beschrieben worden sind, ändern sich allerdings durch den modifizierten Prozess nur minimal. Unterschied dazu ist jetzt, dass die Mitarbeiter*innen drei Kategorisierungsvorschläge bekommen, welche ausgewählt werden können. Trotzdem muss aktuell von den Mitarbeiter*innen jedes Ticket kontrolliert und selbst kategorisiert werden, was den Prozess nicht erleichtert.

Anfangs war es für die künstliche Intelligenz sehr schwierig, aufgrund des Inhalts die passende Kategorie zu finden. Gründe für diese Thematik wurden in diesem Kapitel bereits beschrieben und werden auch durch die Veränderung hinsichtlich künstlicher Intelligenz erneut erläutert. Es ist nach wie vor schwierig, dass die KI richtige Vorschläge für Tickets liefert. Bei einfachen nicht komplexen Tickets sind die Vorschläge meistens sehr gut und können auch ausgewählt werden. Trotzdem muss die vorgeschlagene Kategorisierung nicht genommen werden und die Mitarbeiter*innen haben nach wie vor die Möglichkeit, selbst eine andere Kategorie auszuwählen.

Die Dauer, bis die KI die drei Vorschläge bringt ist nach wie vor zu lange und meist sind die Expert*innen durch ihre Erfahrung schneller in der Kategorisierung als die künstliche Intelligenz. Eine Herausforderung bzw. Schwierigkeit wird bei den Mitarbeiter*innen gesehen, wenn der

Prozess als fixer Bestandteil umgesetzt wird und unumgänglich wird. Hier zweifeln die Mitarbeiter*innen an der Leistung der KI aufgrund von mehreren Faktoren.

Ein Faktor ist beispielsweise die nach wie vor sehr große Anzahl an verfügbaren Kategorien. Die Mitarbeiter*innen selbst wissen teilweise nicht, welche Kategorie zu einhundert Prozent richtig und somit ist es auch für eine KI schwierig, diese Tickets richtig zu kategorisieren. Weiters ist es auch durch die Möglichkeit, dass Tickets nach wie vor ohne Form beispielsweise per Mail, gemeldet werden können schwierig, das Problem aus der Anfrage identifizieren zu können. So kommt es auch nach wie vor aufgrund des ständigen Lernens der KI nach wie vor zu fehlerhaften Kategorisierungen. Trotzdem melden die Mitarbeiter*innen eine Verbesserung der Vorschläge durch das Lernen der zuvor getätigten Kategorien. Hier wird die Anzahl an falsch vorgeschlagenen Kategorien deutlich geringer und ist nur mehr bei komplexen Tickets vorhanden.

Somit sehen die Expert*innen einiges an Potenzial in der Lösung, sehen aber aktuell aufgrund vieler zusammenhängenden Probleme keine Verbesserung im Zeitfaktor bzw. Arbeitsprozess.

4.4.8 Zukunftsaussichten und Entwicklungspotenziale

Im folgenden Kapitel sollen die Erwartungen an zukünftige Entwicklungen im Bereich der künstlichen Intelligenz im Service Desk eingeordnet werden. Es geht dabei auch um mögliche entstehende Diskussionen über Potenziale von verschiedenen Weiterentwicklungen, sowohl im Umfeld der Ticketkategorisierung als auch in anderen Bereichen. Eventuell werden auch diverse Chancen und Risiken bei der Nutzung von KI-Technologien erwähnt bzw. eruiert.

Grundsätzlich haben die Expert*innen durch die Bank erwähnt, dass der Kategorienbaum bzw. auch das Ticketsystem an sich auch schlanker gehalten werden könnte und hier ein Entwicklungspotenzial sichtbar ist. Generell sollte ein Ticket ausnahmslos nach einer vorgefertigten Struktur angelegt werden können, die es den Mitarbeiter*innen erleichtert, den Folgeprozess bis hin zur Lösung des Tickets erledigen zu können. Der Kanal, in welchem das Ticket entsteht, sollte dabei keine Rolle spielen.

Bezüglich der Kategorisierung ist es klar der Wunsch, dass die manuelle Kategorisierung in nächster Zeit komplett wegfallen soll, um einen Mehrwert erzielen zu können. Den Expert*innen ist bewusst, dass dies ein paar Anpassungen voraussetzen würde, wäre aber in gewisser Weise für die Mitarbeiter*innen wünschenswert. Der Betreff könnte als Entwicklungspotenzial auch von der KI generiert werden, um problematische Tickets per Mail mit beispielsweise Betreff „Hilfe!!!“ mit einem richtigen Betreff, welcher auch das Problem an sich betrifft, zu benennen. Weiters kommen auch verschiedene Bestellungen der unterschiedlichen Unternehmen, wie Hardware

oder Benutzereintritte in das Ticket-System und werden immer derselben Kategorie zugewiesen, welche vorab von der künstlichen Intelligenz abgefangen werden könnten.

Abgesehen von der aktuellen Thematik hinsichtlich der Kategorisierung gingen auch andere Potenziale aus den Interviews hervor. Es gibt aktuelle eine eingeschränkte Suche in abgeschlossenen Tickets, welche aber von den Expert*innen für nicht sehr hilfreich empfunden wird. Die Suche von bereits gelösten Tickets ist für einige Expert*innen von Vorteil und hilft bei der Lösungsgenerierung des vorliegenden Problems. Hier sehen die Mitarbeiter*innen auch Potenzial, um die tägliche Arbeit erleichtern zu können. Weiters kam von manchen Expert*innen der Vorschlag, einen Chatbot für die Kund*innen als nächsten Schritt zu implementieren, der die Mitarbeiter*innen für die Fehlerbehandlung unterstützen soll. Aus Sicht der Expert*innen gäbe es vermehrt den Wunsch auch bei den Kund*innen solch eine Lösung zu implementieren, um zu einer schnelleren Lösung des Problems zu gelangen, bevor aufgrund von Überlastungen gewisse Tickets mehrere Tage bis zur Lösung benötigen.

Für zukünftige Implementierungen im Service Desk sehen die Expert*innen einiges an Potenzial in der Abwicklung der Projekte. Es soll in diversen Prozessen eine klare Struktur geschaffen werden, dass alle Mitarbeiter*innen mehr oder weniger an einem Strang ziehen. Weiters sollen die Schulungen bei neu implementierten Veränderungen forciert werden, um den Mitarbeiter*innen ein klares Bild für die Verwendung vermitteln zu können. Künstliche Intelligenz ist aus Sicht der Expert*innen ein Faktor, der in Zukunft nicht wegzudenken ist, sollte aber im Unternehmen auf aktuell eingeführte Prozesse zur Verbesserung angewendet werden. Beispielsweise erwähnten einige Expert*innen, dass diese kein Optimierungspotenzial für eine fünf Sekunden Tätigkeit sehen, da es keine Mitarbeiter*innen gibt, die tagtäglich durchgehend Tickets kategorisieren.

4.4.9 Allgemeine Anmerkungen und Sonstiges

Abschließend werden in dieser Kategorie sonstige allgemeine Anmerkungen angeführt, welche die Expert*innen im Rahmen der Interviews erwähnt haben und über die behandelnde Thematik hinaus gehen. Es wurde damals mit dem Categorizer auch eine weitere Funktion eingeführt, welche „Similar Tickets“ heißt. Diese Funktion stellt den Mitarbeiter*innen drei Tickets zur Verfügung, welche in der Vergangenheit ähnliche Probleme ausgelöst haben und wie sie gelöst worden sind. Somit kann bei Tickets, zu denen aktuell kein Lösungsansatz durch die bearbeitenden Mitarbeiter*innen gefunden werden konnte, ein Lösungsansatz in bereits gelösten Problemen gefunden werden. Das wird von den Expert*innen als hilfreich und positives Feature angesehen.

Allgemein wurde öfters zum Thema künstliche Intelligenz festgehalten, dass dieses Thema im Unternehmen in aktuellen Prozessen nicht einfach ohne jegliche Rücksprache eingeführt werden kann, sondern dies genauer durchdacht gehört und auf die im Unternehmen verwendeten Prozesse angepasst werden müsste, um zum Erfolg zu führen. Die Expert*innen brachten auch mehrmals die Gedanken ein, dass diese Implementierungen sehr viel Zeit und Geld kosten und dann auch sinnvoll implementiert werden sollten.

4.5 Interpretation der Ergebnisse

Zu Beginn der Interpretation ist es wichtig anzumerken, dass die Interpretation der Ergebnisse grundsätzlich mehrere Perspektiven haben kann und somit subjektiv zu betrachten ist. Diese Sichtweise wird durch die Ergebnisse der durchgeführten Interviews sowie eine mehrjährige Erfahrung im Unternehmen des Autors geprägt, durch die auch einige Prozesse detaillierter bekannt sind.

4.5.1 Persönliche Informationen der Expert*innen

Die Analyse der persönlichen Informationen der befragten Expertinnen bietet einen Einblick in das breite Spektrum der Altersgruppen im Service Desk der Styria IT Solutions GmbH. Es ist ermutigend festzustellen, dass Expert*innen aus verschiedenen Altersgruppen vertreten sind, was eine Vielfalt an Erfahrungen und Perspektiven bedeutet. Die Altersgruppen reichen von jungen, unerfahrenen Mitarbeitenden bis hin zu erfahrenen Personen zwischen 50 und 60 Jahren. Die meisten Expert*innen sind zwischen 30 und 39 Jahre alt. Dies zeigt eine gesunde Mischung aus jungen Talenten und erfahrenen Fachleuten.

4.5.2 Vorwissen im Bereich künstliche Intelligenz

Die Befragung zeigt, dass das Interesse und die Einstellung der Expert*innen zu KI stark variieren. Einige haben wenig bis kein Interesse an KI und halten den aktuellen Stand der Technologie für nicht ausgereift genug. Andererseits sehen insbesondere jüngere Expert*innen das Potenzial von KI und nutzen es bereits in verschiedenen Bereichen. Es gibt jedoch eine allgemeine Skepsis und Besorgnis darüber, wie KI in Zukunft eingesetzt werden könnte, besonders in Bezug auf die Grenzen zwischen Realität und künstlicher Intelligenz. Trotzdem wird KI für bestimmte Standardaufgaben wie Bildbearbeitung als nützlich angesehen.

4.5.3 Technische Aspekte der KI-Integration

Ticketsystem allgemein

Die Einschätzungen zum Ticketsystem sind gemischt. Während es als mächtig und komplex wahrgenommen wird, gibt es auch Kritikpunkte, insbesondere in Bezug auf die Benutzerfreundlichkeit und die Vergleichbarkeit mit anderen Systemen. Einige Expert*innen finden das Arbeiten mit Stichwörtern und die vielfältigen Möglichkeiten des Systems nützlich. Jedoch wünschen sich viele eine schlankere Version des Ticketsystems, die besser zur Unternehmensstruktur passt.

Technische Aspekte durch KI

Die Implementierung von KI hat neue technische Aspekte hervorgebracht. Eine zentrale Herausforderung ist die Zeit, die das Tool benötigt, um Kategorisierungsvorschläge zu machen. Die Genauigkeit der vorgeschlagenen Kategorien wird ebenfalls kritisiert, da es häufig zu fehlerhaften oder ungenauen Vorschlägen kommt. Ein weiteres Problem ist die Anpassungsfähigkeit des Systems, da Änderungen im Ticketinhalt nicht immer erkannt werden. Der Kategorisierungsbaum wird als zu komplex empfunden, was zu Verwirrung und Schwierigkeiten bei der richtigen Zuordnung der Tickets führt.

4.5.4 Auswirkungen auf den Arbeitsprozess

Vor KI

Vor der Einführung der KI war der Kategorisierungsprozess nicht klar definiert und es gab keine einheitliche Arbeitsweise unter den Mitarbeiter*innen. Verbesserungsvorschläge wurden nicht immer berücksichtigt oder unzureichend umgesetzt, ohne vorherige Tests. Die Mitarbeiter*innen waren nicht alle über die vom Unternehmen festgelegte Zeitspanne für die Ticketkategorisierung informiert, was zu unterschiedlichen Herangehensweisen führte. Oft wurden Tickets aufgrund von Zeitdruck ohne genaue Überprüfung kategorisiert.

Die Kategorisierungszeit für einfache Tickets lag zwischen fünf und maximal 20 Sekunden pro Ticket. Komplexere Tickets mit längeren Mailverläufen erforderten entsprechend mehr Zeit. Neue Mitarbeiter*innen benötigten ebenfalls mehr Zeit, da sie noch nicht mit den Unternehmensstrukturen vertraut waren.

Nach KI

Nach der Implementierung von KI hat sich der Prozess für die meisten Expert*innen nicht wesentlich geändert. Die KI-gestützten Vorschläge werden oft ignoriert, da viele Mitarbeiterinnen ihre eigenen Erfahrungen und Einschätzungen bevorzugen. Bei einfachen Tickets funktioniert die

KI meistens gut, aber bei komplexen Problemen stößt sie an ihre Grenzen. Die Geschwindigkeit der KI bei der Bereitstellung der Scores wird von vielen Expert*innen als zu langsam empfunden. Es dauert einige Sekunden, bis die Vorschläge erscheinen, währenddessen haben die Mitarbeiter*innen das Ticket bereits selbst kategorisiert. Da die Verwendung der KI-Vorschläge optional ist, verwenden viele Mitarbeiter*innen weiterhin ihre eigenen Einschätzungen.

4.5.5 Benutzerfreundlichkeit und Schulung

Benutzerfreundlichkeit

Die Benutzerfreundlichkeit des KI-unterstützten Systems hat sich verbessert, insbesondere durch kleinere Anpassungen wie die automatische Kategorisierung basierend auf dem Betreff. Allerdings dauert es bei automatisch generierten Tickets, z.B. aus E-Mails, immer noch zu lange für die Kategorisierungsvorschläge, weshalb viele Mitarbeiter*innen weiterhin manuell kategorisieren. Drei Viertel der befragten Expert*innen empfinden die Veränderung im Prozess als nicht besonders relevant, da die manuelle Kategorisierung weiterhin vorherrscht. Wenn eine von der KI vorgeschlagene Kategorie ausgewählt wird, kann diese nicht mehr manuell geändert werden. Die farbliche Markierung der Kategorisierungsvorschläge wird von einigen Expert*innen als hilfreich empfunden, während andere weiterhin manuell kategorisieren.

Schulung

Die fehlende Schulung und Einführung des neuen KI-unterstützten Prozesses wird von allen Expertinnen als kritisch betrachtet. Die Implementierung erfolgte ohne angemessene Unterstützung oder Schulungsmaterialien. Das System wurde einfach live geschaltet, und die Mitarbeiterinnen erhielten nur die Information, dass die Kategorisierung nun durch die KI erfolgen soll. Aufgrund anfänglicher Probleme und Fehlfunktionen des Tools wurde die Einführungsphase als zu kurz und unvollständig empfunden. Eine genauere Testphase wäre für eine effizientere Implementierung erforderlich gewesen.

Insgesamt wird die Nutzung des KI-unterstützten Systems von vielen Expert*innen als nicht unbedingt relevant angesehen, da die manuelle Kategorisierung weiterhin vorherrscht und das Tool als "verfügbar" betrachtet wird, jedoch von den Wenigsten genutzt wird. Die fehlende Schulung hat den Umgang mit dem neuen Tool zusätzlich erschwert.

4.5.6 Mitarbeiterrückmeldung und Akzeptanz

Die Analyse der Mitarbeiterrückmeldungen und der Akzeptanz der künstlichen Intelligenz (KI) im Service Desk zeigt gemischte Reaktionen der Expert*innen.

Feedback

Grundsätzlich kann gesagt werden, dass die Einführung der KI von allen Expert*innen nicht als Verschlechterung angesehen wird, aber auch nicht als klare Verbesserung. Viele Mitarbeiter*innen bevorzugen es, die Kategorisierung aufgrund ihrer eigenen Erfahrung und ihres Wissens vorzunehmen. Sie sind oft schneller und verlassen sich nicht blind auf die Vorschläge der KI. Die Mitarbeiter*innen würden eine schnellere Anzeige der Kategorisierungsvorschläge begrüßen. Allerdings müssen sie auch weiterhin sorgfältig überprüfen, ob der Vorschlag korrekt ist. Die Implementierung der KI wird von neuen Mitarbeiter*innen positiver wahrgenommen, da die Vorschläge der KI ihnen eine Richtung geben können und die Einarbeitung in die Materie erleichtern.

Akzeptanz

Zu Beginn der Implementierung und ohne einen klaren Einführungs- oder Schulungsprozess waren viele Expert*innen skeptisch und negativ gestimmt gegenüber der KI. Da die Verwendung der KI nicht verpflichtend ist und viele Expert*innen weiterhin ihren eigenen Prozess bevorzugen, hat die KI keine großen Auswirkungen auf den Arbeitsprozess. Die anfänglich negative Stimmung hat sich im Laufe der Zeit verbessert. Einige Expert*innen, die anfangs gegen die Implementierung waren, sind mittlerweile neutraler gestimmt und probieren die Kategorisierung ab und zu aus. Einige Expert*innen kritisieren, dass die Implementierung in der aktuellen Form unnötig ist und keine wirkliche Verbesserung im Prozess bietet. Sie sehen das investierte Geld und die Zeit als potenziell anders investierbar.

4.5.7 Herausforderungen und Schwierigkeiten

Die Analyse zeigt, dass vor der Implementierung der künstlichen Intelligenz im Service Desk bereits einige Herausforderungen und Schwierigkeiten bei der Ticketkategorisierung existierten. Eine der Hauptprobleme war die richtige Zuordnung der Tickets zu den entsprechenden Teams. Falsche Kategorisierungen führten zu zusätzlichem Aufwand, da die Abteilungen die Tickets prüfen mussten, auch wenn sie nicht dafür zuständig waren. Neue Mitarbeiter*innen hatten Schwierigkeiten, die Vielzahl an Kategorien richtig zu verwenden und die zuständige Abteilung für ein bestimmtes Problem zu identifizieren. Oft waren mehrere Kategorien für ein Ticket geeignet, aber hinter den Kategorien standen verschiedene Abteilungen, was die Auswahl erschwerte.

Ein weiteres Problem war, dass Kunden oft mehrere Probleme in einem Ticket meldeten, was die Kategorisierung komplizierter machte. Die KI sollte dazu dienen, einige dieser Probleme zu lösen. Nach der Implementierung der KI erhielten die Mitarbeiter*innen drei Kategorisierungsvorschläge, mussten jedoch weiterhin jedes Ticket selbst kontrollieren und kategorisieren.

Die Analyse zeigt, dass die KI anfangs Schwierigkeiten hatte, die richtige Kategorie aufgrund des Inhalts des Tickets zu finden. Die Qualität der Kategorisierungsvorschläge variierte je nach Komplexität des Tickets. Die Mitarbeiter*innen haben immer noch die Möglichkeit, eine andere Kategorie auszuwählen, wenn sie den Vorschlag der KI nicht akzeptierten.

4.5.8 Zukunftsaussichten und Entwicklungspotenziale

Die Expert*innen haben deutliche Erwartungen an zukünftige Entwicklungen im Bereich der künstlichen Intelligenz im Service Desk. Eine zentrale Anforderung ist die Überarbeitung des Kategorienbaums und des Ticketsystems, um sie schlanker und effizienter zu gestalten. Dies soll es den Mitarbeiter*innen erleichtern, Tickets nach einer vorgefertigten Struktur anzulegen, unabhängig vom Kanal, über den das Ticket eingereicht wird.

Ein entscheidender Wunsch der Expert*innen ist der Wegfall der manuellen Kategorisierung zugunsten einer automatisierten Lösung durch künstliche Intelligenz. Diese Änderung erfordert Anpassungen, wird aber als bedeutender Mehrwert für die Mitarbeiter*innen betrachtet. Zusätzlich wird vorgeschlagen, dass die KI den Betreff von eingehenden Mails generieren könnte, um problematische Tickets besser zu benennen.

Des Weiteren gibt es Potenzial für die KI, wiederkehrende Bestellungen wie Hardware oder Benutzereintritte automatisch einer spezifischen Kategorie zuzuweisen. Dies würde Zeit sparen und den Arbeitsprozess optimieren.

Abseits der Kategorisierung wurde auch die eingeschränkte Suchfunktion in abgeschlossenen Tickets als problematisch identifiziert. Die Expert*innen sehen hier Verbesserungspotenzial, da eine effiziente Suche nach bereits gelösten Problemen die Arbeitsprozesse beschleunigen könnte.

Einige Expert*innen schlagen außerdem die Implementierung eines Chatbots für Kund*innen vor, um die Mitarbeiter*innen bei der Fehlerbehandlung zu unterstützen. Dies könnte zu einer schnelleren Lösung von Problemen führen und Überlastungen bei Tickets reduzieren.

Für zukünftige Implementierungen im Service Desk sehen die Expert*innen Potenzial in der Schaffung klarer Prozessstrukturen, um eine einheitliche Arbeitsweise zu gewährleisten. Zusätzlich sollte die Schulung der Mitarbeiter*innen bei neuen Veränderungen intensiviert werden, um das Verständnis und die Effizienz zu verbessern. Künstliche Intelligenz wird als unverzichtbarer Faktor für die Zukunft des Unternehmens angesehen, sollte jedoch in bestehende Prozesse integriert werden, um sie zu verbessern.

4.5.9 Allgemeine Anmerkungen und Sonstiges

Die Expert*innen haben in ihren Interviews einige wichtige Punkte hervorgehoben, die über die unmittelbare Thematik der Ticketkategorisierung hinausgehen. Ein solcher Punkt ist die Einführung der Funktion "Similar Tickets" durch den Categorizer. Diese Funktion stellt den Mitarbeiter*innen drei Tickets zur Verfügung, die in der Vergangenheit ähnliche Probleme verursacht haben und wie sie damals gelöst wurden. Dieses Feature wird von den Expert*innen als hilfreich und positiv bewertet, da es die Mitarbeiter*innen bei der Lösungsfindung unterstützt, insbesondere wenn für aktuelle Tickets noch kein Lösungsansatz gefunden wurde.

Ein weiterer wichtiger Aspekt, der von den Expert*innen hervorgehoben wurde, betrifft die Einführung von künstlicher Intelligenz im Unternehmen. Es wird betont, dass die Einführung von KI nicht einfach ohne Rücksprache oder Planung in bestehende Prozesse integriert werden sollte. Vielmehr muss dieser Schritt sorgfältig durchdacht und auf die spezifischen Prozesse und Bedürfnisse des Unternehmens angepasst werden, um erfolgreich zu sein. Die Expert*innen betonen auch, dass solche Implementierungen sowohl Zeit als auch Geld erfordern und daher sinnvoll und effektiv umgesetzt werden müssen.

4.6 Handlungsempfehlung

Die vorangegangene Analyse hat die verschiedenen Aspekte des Kategorisierungsprozesses im Service Desk mit künstlicher Intelligenz beleuchtet. Sowohl vor als auch nach der Implementierung wurden Herausforderungen, Erfahrungen und Auswirkungen auf den Arbeitsprozess der Mitarbeiter*innen untersucht. Die Erkenntnisse zeigen, dass die aktuelle Lösung durch die KI zwar Potenzial hat, aber auch weiterhin auf einige Schwierigkeiten und Herausforderungen stößt.

In diesem Kapitel der Handlungsempfehlung werden konkrete Maßnahmen und Schritte vorgeschlagen, um den Kategorisierungsprozess im Service Desk zu optimieren. Basierend auf den identifizierten Problemen und den Feedbacks der Expert*innen werden Empfehlungen formuliert, die darauf abzielen, die Effizienz, Genauigkeit und Benutzerfreundlichkeit des KI-unterstützten Systems zu verbessern.

4.6.1 Schulung und Einführung

Es sollten im Unternehmen gezielte Schulungen und dadurch auch umfassende Schulungsmaterialien entwickelt werden, welche auf verschiedene Altersgruppen und Erfahrungsgruppen der Expert*innen angewendet werden können. Die Schulungen sollten dabei nicht nur die technischen Aspekte der KI-Integration abdecken, sondern auch die Vorteile für den

zukünftigen Arbeitsprozess verdeutlichen. Weiters sollten die Mitarbeiter*innen regelmäßig über die Verwendung der KI informiert und auch geschult werden, denn regelmäßige Schulungen sind unerlässlich, um sie mit neuen Prozessen und Tools bzw. Veränderungen in diesen vertraut zu machen. Es soll dabei eine klare Struktur für die Prozesse geschaffen werden, um die Arbeitsweise im Service Desk vereinheitlichen zu können.

Der Fokus für die Nutzung von KI sollte darauf liegen, wie effektiv künstliche Intelligenz genutzt werden kann und es soll aufgezeigt werden, wie sich dadurch die Arbeitsprozesse verbessern können. Im konkreten Beispiel soll in der Ticketkategorisierung gezeigt werden, wie Zeit gespart werden kann und Fehler minimiert werden. Ein konkreter Einführungsplan mit klarer Struktur kann dabei hilfreich sein. Dieser beinhaltet Testphasen, in denen Feedback gesammelt werden kann, um das System in weiterer Folge dementsprechend anpassen zu können.

4.6.2 Anpassungen im Ticketsystem

In einem ersten Schritt sollte der aktuell verwendete Kategorienbaum überarbeitet werden. Durch eine Vereinfachung des Kategorienbaums kann die Verwirrung bei der Kategorisierung reduziert werden und die Genauigkeit bei der Ticketzuweisung verbessern. Weiters ermöglicht ein vereinfachter Kategorienbaum eine schnellere und genauere Kategorisierung. Somit kann es hilfreich sein, die Anzahl der verfügbaren Kategorien zu reduzieren und eine klare Richtlinie für die Auswahl der Kategorien zu erstellen. Dieser Schritt kann in erster Linie die Mitarbeiter*innen entlasten und in weiterer Folge auch die KI-Genauigkeit verbessern. Die durch künstliche Intelligenz unterstützte Lösung in der Ticketkategorisierung sollte kontinuierlich trainiert und verbessert werden, um genauere Kategorisierungsvorschläge zu liefern. Dies kann durch Überwachung und regelmäßigen Updates erfolgen.

Im nächsten Schritt sollte eine Automatisierung der Kategorisierung durch die KI erfolgen, um die manuelle Arbeit zu minimieren. Um die Umstellung aus Angst vor Fehlern nicht in einem großen Schritt machen zu müssen, bietet es sich an, die Automatisierung bei sehr einfachen Tickets zu starten, welche immer dieselbe Kategorie zugewiesen bekommen. Dies wäre für den Start ein erster Schritt und kann den Mitarbeiter*innen Zeit sparen und die Effizienz steigern. Anpassungen im System und Schulungen für Mitarbeiter*innen sind erforderlich, um diesen Übergang reibungslos zu gestalten. Die KI könnte weiters genutzt werden, um die Betreffzeilen problematischer Tickets zu generieren, was die Identifizierung und Bearbeitung erleichtern würden.

Als weitere Anpassung wird noch eine klare Strukturierung von Tickets in der Erstellung empfohlen. Kund*innen sollten angeleitet sein, klare und strukturierte Informationen in ihren

Tickets bereitzustellen. Dieser Schritt könnte mit verbesserter Kommunikation und mit Richtlinien für die Ticketerstellung erreicht werden.

4.6.3 Verbesserung der Benutzerfreundlichkeit

Die Benutzerfreundlichkeit des Systems ist ein nächster Punkt, der verbessert werden kann. Begonnen wird mit der Empfehlung, die Anzeige von Kategorisierungsvorschlägen zu beschleunigen. Die Geschwindigkeit sollte optimiert werden, um sicherzustellen, dass die Vorschläge noch relevant sind, wenn die Mitarbeiter*innen die Tickets bearbeiten. Es sollte dabei eruiert werden, ob die Kategorisierung von der KI bereits vorab berechnet werden kann, dass bei erstmaligem Öffnen der Mitarbeiter*innen die Vorschläge bereits präsent sind. Weiters soll es den Expert*innen erlaubt sein, auch nach einer getätigten Auswahl eines Vorschlages, die Kategorie manuell anzupassen. Dieser Schritt würde das Vertrauen in das System fördern und ermöglicht gleichzeitig eine präzisere Zuweisung der Tickets. Im Zusammenhang mit der Geschwindigkeit sollte auch bei Änderungen im Betreff oder im Ticketinhalt, die bereits vorgeschlagene Kategorie automatisch angepasst werden können.

Das Ticketportal selbst gilt es auch benutzerfreundlicher zu gestalten, um den Nutzen der Kund*innen zu erhöhen. Es könnten dadurch auch strukturierte Tickets erstellt werden, was die Kategorisierung erleichtern würde. Der Prozess, dass ungefilterter Inhalt per Mail an den Service Desk gesendet werden kann, müsste ebenfalls mit Hilfe eines Formulars angepasst werden, um eine Struktur implementieren zu können. Als weiteren Schritt kann in Betracht gezogen werden, einen Chatbot zu implementieren, um die Mitarbeiter*innen bei der Fehlerbehebung zu unterstützen. Der Chatbot kann zur schnelleren Lösung von Problemen führen und die Arbeitslast der Mitarbeiter*innen reduzieren. Hier kann die Implementierung auch in mehreren Schritten erfolgen und zuerst ein Chatbot für die Expert*innen intern entwickelt werden, der in internen Dokumentationen und Lösungen im Ticketsystem sucht und Lösungsvorschläge zurücksendet. In einem weiteren Schritt könnte dann die Interaktion mit den Kund*innen, wie mit einem klassischen Chatbot, forciert werden.

4.6.4 Kommunikation und Feedback

Nicht nur hinsichtlich der Entwicklung von künstlicher Intelligenz, sondern auch generell sollten regelmäßige Feedback-Meetings einberufen werden, um die Meinungen und Erfahrungen der Expert*innen sammeln zu können. Dies ermöglicht hinsichtlich der KI eine kontinuierliche Verbesserung der KI-Integration, kann aber auch generell für eine konstruktive Kommunikation auf alle Prozesse ausgeweitet werden. Weiters kann auch ein Feedbackmechanismus entwickelt werden, welcher es ermöglicht, Verbesserungsvorschläge und Probleme direkt an das

Entwicklungsteam des KI-Systems zu kommunizieren. Dadurch könnten Anpassungen vorgenommen werden, die den Arbeitsprozess effizienter und benutzerfreundlicher machen. Das Feedback sollte aktiv bei den Expert*innen eingeholt und berücksichtigt bzw. diskutiert werden. Es sollte auch stets darauf geachtet werden, dass die Vorteile der KI für den Service Desk deutlich kommuniziert werden, insbesondere die effizienteren und präziseren Arbeitsprozesse und Arbeitsabläufe sollten den Mitarbeiter*innen präsentiert werden.

4.6.5 Zukünftige Entwicklungen

Wie bereits kurz davor erwähnt, sollte die automatische Zuweisung von wiederkehrenden Tickets forciert werden. Darunter können unter anderem Hardware-Bestellungen oder auch Benutzereintritte als Beispiel dienen und automatisiert von der KI kategorisiert werden. Dies würde erheblich Zeit sparen und optimiert den Arbeitsprozess.

Die Suchfunktion für abgeschlossene Tickets sollte weiterentwickelt und verbessert werden, um eine effizientere Suche nach bereits gelösten Problemen zu ermöglichen. Das Ticketsystem selbst sollte eine effiziente Suchfunktion bieten, um die Lösungsgenerierung und die Mitarbeiter*innen zu unterstützen. Dies könnte durch die Integration von weiteren Suchalgorithmen erreicht werden.

Wie bereits von den Expert*innen in den Interviews erwähnt, ist die „Similar-Tickets-Funktion“ für einige sehr hilfreich in der Fehleranalyse und sollte regelmäßig überprüft werden, um sicherzustellen, dass es die Mitarbeiter*innen weiterhin bei der Lösungsfindung unterstützt. Auch gilt es zu eruieren, wie das Feature entwickelt werden kann, um alle Mitarbeiter*innen zu überzeugen, es zu nutzen.

4.6.6 Fortlaufende Evaluierungen

Zukünftig sollten regelmäßige Überprüfungen der KI-Integration eingeplant werden, um sicherzustellen, dass diese Implementierungen den Bedürfnissen des Service Desks entsprechen. Es soll dabei keine Barriere entstehen, Systemanpassungen aufgrund des Feedbacks der Mitarbeiter*innen zu tätigen und den sich in der Zukunft entwickelnden Anforderungen.

Aus Sicht der Kosten sollte das Unternehmen eine gründliche Überprüfung des Return on Investment (ROI) der KI-Implementierung durchführen. Sollte das System trotz mehreren Anpassungen und Entwicklungen nicht die erwarteten Verbesserungen oder Effizienzgewinne bringen, sollten alternative Investitionen in Betracht gezogen werden, die möglicherweise einen größeren Nutzen für den Service Desk bieten können.

Abschließend kann gesagt werden, dass es wichtig ist, eine offene Einstellung gegenüber neuen Technologien und deren Implementierungen zu haben. Eine offene Unternehmenskultur ist entscheidend für den Erfolg von diversen KI-Integrationen. Durch die Umsetzung dieser Handlungsempfehlungen könnte die KI-Integration und die generelle Struktur im Service Desk optimiert werden. Dies würde die vorhin aufgezeigten Probleme beseitigen, zu einem effizienteren Arbeitsprozess bzw. einer präziseren Ticketzuweisung und einer verbesserten Zufriedenheit der Mitarbeiter*innen beitragen.

5 CONCLUSIO

Abschließend wird in Kapitel 5 das Thema Conclusio behandelt. Begonnen wird mit einer Zusammenfassung bzw. einer Schlussfolgerung, inklusive Beantwortung der Forschungsfrage. Abgeschlossen wird der empirische Teil mit einem Ausblick für zukünftige Forschungsthemen in diesem Bereich.

5.1 Schlussfolgerung

Folgende Forschungsfrage lag der vorliegenden Masterarbeit zu Grunde:

„Wie beeinflusst die Verwendung von künstlicher Intelligenz die Effizienz und Arbeitsbelastung des Supportteams in einem Ticketsystem?“

Die vorliegende Arbeit hat somit zum Ziel, die Beeinflussung von künstlicher Intelligenz auf Effizienz und Arbeitsbelastung der Mitarbeiter*innen aufzuzeigen. Aufgrund der durchgeführten Experteninterviews kann in der aktuellen Form der Implementierung für die Styria IT Solutions folgendes gesagt werden. Die Implementierung von künstlicher Intelligenz in verschiedenen Arbeitsprozessen ist ein weiterer Schritt in Richtung Automatisierung von Prozessen. Die Analyse der Interviews gibt einen tiefen Einblick in die Auswirkungen auf das Supportteam und die Effizienz des Ticketsystems. Es ist klar ersichtlich, dass das Interesse und die Einstellung gegenüber künstlicher Intelligenz bei den Expert*innen stark variiert, wobei einige Mitarbeiter*innen großes Potenzial darin erkennen und andere eher skeptisch sind. Dies kann auf die verschiedenen Altersgruppen und der breiten Palette an Perspektiven und Erfahrungen zurückzuführen sein.

Bezüglich der technischen Aspekte der KI-Integration wurden Herausforderungen identifiziert, welche beispielsweise längere Zeit für Kategorisierungsvorschläge oder auch ungenaue Kategorisierungsvorschläge sind. Die komplexe Struktur des aktuell verwendeten Kategorienbaums ist eine weitere Herausforderung. Der Kategorisierungsprozess ist aktuell nicht einheitlich und führt zu inkonsistenten Ergebnissen.

Die Arbeitsweise hat sich bei vielen Expert*innen auch mit der Einführung von künstlicher Intelligenz nicht wesentlich verändert. Viele nutzen die KI-Vorschläge nicht aktiv oder benutzen ihre eigenen Einschätzungen aufgrund jahrelanger Erfahrung. Diese Punkte deuten darauf hin, dass die Akzeptanz und Nutzung der KI nicht optimal sind.

Somit kann zur Beantwortung der Forschungsfrage folgendes gesagt werden. Die Verwendung von künstlicher Intelligenz im Ticketsystem des Service Desks der Styria IT Solutions hat einige Herausforderungen hinsichtlich Effizienz und Arbeitsbelastung zu bewältigen. Diese wären unter

anderem die langsame Geschwindigkeit der KI, bis Vorschläge aufgrund des Inhalts generiert werden. Die manuellen Präferenzen der Expert*innen sind ebenfalls eine Herausforderung, welche zu Inkonsistenzen in der Kategorisierung führen und eine Vereinheitlichung erschweren. Die komplexe Struktur des Kategorienbaums führt oft zu Verwirrung und kann somit zu Fehlkategorisierungen führen und die Arbeitslast damit erhöhen. Positiv kann trotzdem erwähnt werden, dass die Ticketkategorisierung durch künstliche Intelligenz eine richtungsweisende Unterstützung für neue Mitarbeiter*innen sein kann und bei der Einarbeitung in den Kategorisierungsprozess unterstützen kann.

Insgesamt kann festgestellt werden, dass die Verwendung von künstlicher Intelligenz im Ticketsystem des Service Desks potenziell die Effizienz steigern kann, insbesondere bei der Bearbeitung einfacher Tickets und der Unterstützung neuer Mitarbeiter*innen. Jedoch müssen noch Herausforderungen überwunden werden, wie die langsame Geschwindigkeit der KI, die manuellen Präferenzen der Expert*innen und die Komplexität des Kategorienbaums. Eine gezielte Schulung, kontinuierliches Feedback und mögliche Anpassungen am System können dazu beitragen, diese Herausforderungen zu bewältigen und die Effizienz des Supportteams langfristig zu verbessern. Die Implementierung von künstlicher Intelligenz im Service Desk bietet somit Potenzial für eine effizientere Ticketbearbeitung und eine bessere Auslastung der Ressourcen, wenn die vorhandenen Herausforderungen erfolgreich angegangen werden.

5.2 Ausblick

Die Ergebnisse lassen darauf schließen, dass es im Bereich künstliche Intelligenz viele Forschungsmöglichkeiten gibt. Für zukünftige Forschungen bieten sich verschiedene Bereiche an, welche vertieft werden könnten, um weitere Optimierungen vornehmen zu können. Einige Vorschläge werden folgend aufgelistet:

- Evaluierung verschiedener KI-Modelle und -Technologien, um die Effizienz und Effektivität zu maximieren
- Erforschung der Implementierung von Chatbots zur direkten Unterstützung von Kund*innen und zur Entlastung des Supportteams
- Methodenforschung für eine effektivere Zusammenarbeit zwischen Expert*innen und KI
- Untersuchung einer optimalen Schulungsstrategie für die Expert*innen, um eine bessere Akzeptanz und Nutzung der KI zu fördern
- Vergleichsstudien mit anderen Unternehmen, welche die identen Lösungen implementieren, um bewährte Verfahren und Erfolgsmuster bei der KI-Integration zu identifizieren

- Evaluierung von unterschiedlichen KI-Systemen, um die am besten geeigneten Lösungen für den Service Desk zu ermitteln
- Erforschung der Anwendbarkeit von KI in anderen Bereichen im Service Desk, wie zum Beispiel in der Problemlösung oder Kund*innenkommunikation

Insgesamt bieten diese möglichen Forschungsbereiche einen umfassenden Ausblick für weitere Untersuchungen zur Integration von künstlicher Intelligenz im Service Desk. Durch vertiefte Analysen, Evaluierungen und Implementierungen können zukünftige Studien dazu beitragen, die Effizienz, Arbeitsbelastung und Qualität des Supports weiter zu verbessern und die Potenziale der KI voll auszuschöpfen.

ANHANG A - Interviewleitfaden

Persönliche Informationen

- 1.) Wie alt sind Sie?
- 2.) Wie lange arbeiten Sie bereits im Bereich Service Desk?
- 3.) Wie lange arbeiten Sie bereits in der Styria IT Solutions?

Technische Informationen

- 4.) Welche Erfahrungen haben Sie mit dem aktuellen Ticketsystem "Matrix 42"?
- 5.) Welche Erfahrungen haben Sie mit künstlicher Intelligenz? (Beruflich und Privat)

Befragung vor der Implementierung des "Categorizers"

- 6.) Wie ist die aktuelle Ist-Situation im Bearbeitungsprozess der Tickets im First-Level-Support?
- 7.) Wie lange brauchen Sie im Schnitt für die Kategorisierung eines Tickets?
- 8.) Welche Herausforderungen gibt es bei der Kategorisierung der Tickets?
- 9.) Welches Potenzial sehen Sie im aktuellen Bearbeitungsprozess? Wie kann der Prozess verbessert werden, um Ihre Arbeit effizienter gestalten zu können?
- 10.) Mögliche weitere Fragen aufgrund der Gesprächssituation

Befragung nach der Implementierung des "Categorizers"

- 11.) Beschreiben Sie bitte Ihre Zusammenarbeit mit dem KI-System bei der Ticketkategorisierung. Wie ist die neue Situation im Bearbeitungsprozess der Tickets im First-Level Support?
- 12.) Wie lange brauchen Sie im Schnitt für die Kategorisierung eines Tickets?
- 13.) Welche Herausforderungen gibt es bei der Kategorisierung der Tickets?
- 14.) Welche Faktoren beeinflussen Ihre Entscheidungen bei der Auswahl eines bestimmten Scores?
- 15.) Wie bewerten Sie die Genauigkeit, des von der KI vorgeschlagenen Scores im Vergleich zur manuellen Kategorisierung?
- 16.) Haben Sie Situationen erlebt, in denen die KI-basierende Kategorisierung fehlerhaft war? Wenn ja, könnten Sie Beispiele nennen?
- 17.) Hat sich die Situation Ihrer Ansicht nach durch die künstliche Intelligenz verbessert oder verschlechtert?
 - a. Was hat sich konkret am neuen Prozess verbessert?
 - b. Was ist konkret schlechter als zuvor? Sehen Sie persönlich einen Nachteil in der Arbeitsweise für den Arbeitsalltag?
- 18.) Inwieweit hat die Einführung der Kategorisierungsvorschläge Ihre tägliche Arbeit beeinflusst?
- 19.) Wie ist die generelle Akzeptanz des KI-gesteuerten Ticketsystems unter den Mitarbeiter*innen im Service Desk?
- 20.) Gab es Widerstände gegen die Implementierung? Wenn ja, warum?

Ausblick

- 21.) Welche Entwicklungen erwarten Sie im Bezug auf die Ticketkategorisierung im Service Desk?
- 22.) Welche Herausforderungen oder Schwierigkeiten sehen Sie in der aktuellen Implementierung?
- 23.) Gibt es sonstige Anmerkungen oder Informationen, die Sie im Bezug auf die Kategorisierung teilen möchten?

ANHANG B - Interview-Vereinbarung

Einverständniserklärung zur Erhebung und Verarbeitung von Interviewdaten.

Interviewte Person:

Datum:

Ort:

Ich erkläre mich mit meiner Unterschrift dazu bereit, im Rahmen der Masterarbeit „Implementierung von künstlicher Intelligenz im Bearbeitungsprozess von Serviceanfragen im Kundensupport am Beispiel der Styria IT Solutions GmbH & Co KG“ an einem Interview teilzunehmen. Ich wurde über Art, Umfang und Ziel sowie den Verlauf des Interviews informiert.

Dazu bin ich auch damit einverstanden, dass die oben genannte Forschung aufgezeichnet und anschließend in schriftliche Form gebracht und analysiert werden darf. Für die weitere wissenschaftliche Auswertung des Interviewtextes werden alle Angaben, die zu einer Identifizierung der Person oder von im Interview erwähnten Personen und Institutionen führen könnten, anonymisiert. Das Transkript des Interviews dient nur zu Analysezwecken und wird rein in Ausschnitten zitiert.

Die Teilnahme an der Erhebung und die Zustimmung zur Verwendung der Daten sind freiwillig. Durch die Ablehnung entstehen mir keine Nachteile. Unter diesen Bedingungen erkläre ich mich bereit, das Interview zu geben und bin damit einverstanden, dass es aufgezeichnet, verschriftlicht, anonymisiert und ausgewertet werden.

Unterschrift Interviewte*r

Unterschrift Interviewer

ABKÜRZUNGSVERZEICHNIS

CC – Competence Center
IT – Information Technology
DSR – Design Science Research
IS - Informationssysteme
ITSM – IT Service Management
ITIL – Information Technologie Infrastructure Library
SVS – Service Value System
MOF – Microsoft Operations Framework
KPI – Key Performance Indicator
KI – Künstliche Intelligenz
SVM – Support Vector Maschine
KNN – K-nearest Neighbor
RF – random forrest
LR – logistische Regression
NLP – Natural Language Processing
NN – Nearest neighbor rule

ABBILDUNGSVERZEICHNIS

Abbildung 1: Organigramm der Styria IT Solutions GmbH & Co KG (in Anlehnung an (Styria IT, 2023)	1
Abbildung 2: DSR-Ansatz nach Hevner (in Anlehnung an (Peffers et al., 2007))	5
Abbildung 3: Prozess der gestaltungsorientierten Wirtschaftsinformatik (eigene Darstellung, 2023)	7
Abbildung 4: Grundsätze ordnungsmäßiger Modellierung (eigene Darstellung, 2023)	8
Abbildung 5: Empirische Evaluation von Artefakten (in Anlehnung an Wilde & Hess, 2023)	9
Abbildung 6: ITIL Service Value System (SVS) (in Anlehnung an Randus IT, 2023)	12
Abbildung 7: Prozess supervised learning (in Anlehnung an Liu & Wu, 2023)	21
Abbildung 8: Vergleich KNN bei K=1 und K=4 (in Anlehnung an Imanodoust & Bolandraftar, 2023)	24
Abbildung 9: Unklassifiziertes Ticket aus dem Ticketsystem mit Categorizer (eigene Darstellung, 2024) ..	27
Abbildung 10: Klassifiziertes Ticket aus dem Ticketsystem mit Categorizer (eigene Darstellung, 2024)..	27
Abbildung 11: Qualitative Inhaltsanalyse nach Kuckartz (in Anlehnung an Kuckartz & Rädiker ,2022)	36
Abbildung 12: Kategorisierung (eigene Darstellung, 2024)	41
Abbildung 13: Erfahrung im Service Desk nach Altersgruppen (eigene Darstellung, 2024)	42

TABELLENVERZEICHNIS

Tabelle 1 : Unterschiede von KNN-Varianten (eigene Darstellung, 2023).....	23
Tabelle 2 - Ablauf von qualitativen Interviews (eigene Darstellung, 2024)	31

LITERATURVERZEICHNIS

- Abdalla, H. B. (2022). A brief survey on big data: technologies, terminologies and data-intensive applications. *Journal of Big Data*, 9(1). <https://doi.org/10.1186/s40537-022-00659-3>
- Adyrbai, D. (2021). *ITIL Practices in 2000 words: Incident management, service desk and service request management*, Axelos. Verfügbar unter: <https://www.axelos.com/resource-hub/white-paper/itil-2000-words-incident-service-desk-request-management#:~:text=Definition%3A%20Incident&text=of%20a%20service.-,The%20incident%20management%20practice%20ensures%20that%20periods%20of%20unplanned%20service,quick%20restoration%20of%20normal%20operation.>
- Ahmad, Norita; Shamsudin, Zulkifli M. (2013): Systematic Approach to Successful Implementation of ITIL. In: *Procedia Computer Science* 17, S. 237–244. DOI: 10.1016/j.procs.2013.05.032
- Becker, J. (Ed.). (2009a). *Wissenschaftstheorie und gestaltungsorientierte Wirtschaftsinformatik* (1. Aufl.). Heidelberg: Physica-Verl. Abgerufen von: http://bvbr.bib-bvb.de/8991/F?func=service&doc_library=BVB01&doc_number=017602538&line_number=0002&func_code=DB_RECORDS&service_type=MEDIA
- Beims, M. & Ziegenbein, M. (2023). *IT-Service-Management in der Praxis mit ITIL®. Zusammenarbeit systematisieren und relevante Ergebnisse erzielen* (6., aktualisierte Auflage). München: Hanser.
- Brewster, E., et. al. (2016). *IT Service Management::Support for Your ITSM Foundation Exam* (3. Aufl.). Swindon: BCS. <https://web-p-ebscohost-com.elibrary-campus02.at/ehost/detail/detail?vid=0&sid=ae3d1347-198c-432d-ad2b-838a0e238512%40redis&bdata=JkF1dGhUeXBIPWNvb2tpZSxpcCx1cmwsdWlkLHNoaWlmc2l0ZT1laG9zdC1saXZl#AN=1163872&db=e020mww>
- Chen, S. (2018). K-Nearest Neighbor Algorithm Optimization in Text Categorization. *IOP Conference Series: Earth and Environmental Science*, 108, 52074. <https://doi.org/10.1088/1755-1315/108/5/052074>
- Deistler, N. & Rentrop, C. (2022). IT-Compliance in KMU – Experteninterviews zum Status quo. *Wirtschaftsinformatik & Management*, 14(1), 10–19. <https://doi.org/10.1365/s35764-021-00380-5>
- Döring, N. (2023). *Forschungsmethoden und Evaluation in den Sozial- und Humanwissenschaften* (Springer-Lehrbuch, 6., vollständig überarbeitete, aktualisierte und erweiterte Auflage). Berlin, Heidelberg: Springer. <https://doi.org/10.1007/978-3-662-64762-2>

- Fischer, A. (2020). *Von Augmented Reality bis KI. Die wichtigsten IT-Themen, die Sie für Ihr Unternehmen kennen müssen*. München: Hanser.
- Frick, N., Brünker, F., Ross, B. & Stieglitz, S. (2019). Der Einsatz von künstlicher Intelligenz zur Verbesserung des Incident Managements. *HMD Praxis der Wirtschaftsinformatik*, 56(2), 357–369. <https://doi.org/10.1365/s40702-019-00505-w>
- Gabler Wirtschaftslexikon. (2021, 7. Juni). *Definition Big Data*. Verfügbar unter: <https://wirtschaftslexikon.gabler.de/definition/big-data-54101/version-384381>
- Gericke, A. & Robert, W. (2009). Entwicklung eines Bezugsrahmens für Konstruktionsforschung und Artefaktkonstruktion in der gestaltungsorientierten Wirtschaftsinformatik. In J. Becker (ed.), *Wissenschaftstheorie und gestaltungsorientierte Wirtschaftsinformatik* (1. Aufl., S. 195–210). Heidelberg: Physica-Verl. https://doi.org/10.1007/978-3-7908-2336-3_10
- Güttinger, A. (17.05.2023). Persönliche Kommunikation zum Thema Implementierung von künstlicher Intelligenz und Prozessstruktur im Service Desk der Styria IT. (B. Glanznig, Interviewer)
- Imandoust, S. B. & Bolandraftar, M. (2013). Application of K-nearest neighbor (KNN) approach for predicting economic events theoretical background. *Int J Eng Res Appl*, 3, 605–610.
- Jörg Becker, Roland Holten, Ralf Knackstedt & Björn Niehaves. (2003). *Forschungsmethodische Positionierung in der Wirtschaftsinformatik: Epistemologische, ontologische und linguistische Leitfragen*. Arbeitsberichte des Instituts für Wirtschaftsinformatik (93). Münster. Verfügbar unter: <http://hdl.handle.net/10419/59562>
- Jung, R. (Hrsg.). (2008). *Quo vadis Wirtschaftsinformatik? Festschrift für Prof. Gerhard F. Knolmayer zum 60. Geburtstag* (Gabler Edition Wissenschaft, 1. Aufl.). Wiesbaden: Gabler.
- Know-Center. (2023). *AI - state-of-the-art*, Research Center for Data-Driven Business & Big Data Analytics. Verfügbar unter: <https://www.know-center.at/en/ai-state-of-the-art/>
- Kuckartz, Udo; Rädiker, Stefan (2022): *Qualitative Inhaltsanalyse. Methoden, Praxis, Computerunterstützung. Grundlagentexte Methoden*. 5. Auflage. Weinheim, Basel: Beltz Juventa (Grundlagentexte Methoden).
- Kloiber, Thomas (03.10.2023). Persönliches Gespräch zum Thema Categorizer der Firma Leftshift One. (B. Glanznig, Interviewer)
- Liu, Q. & Wu, Y. (2012). Supervised Learning. In N. M. Seel (Hrsg.), *Encyclopedia of the Sciences of Learning* (S. 3243–3245). Boston, MA: Springer US. https://doi.org/10.1007/978-1-4419-1428-6_451

- Logan, R. K. & Braga, A. (Eds.). (2020). AI and the singularity. A fallacy or a great opportunity? Basel, Switzerland: MDPI - Multidisciplinary Digital Publishing Institute. <https://doi.org/69006>
- Marx, J. P. S. (2023). *shrike!* Verfügbar unter: <https://shrike.de/design-science-research-methodologie/>
- McCarthy, J., Rochester, N., Shannon, C. & Minsky, M. (1955). *A proposal for the dartmouth summer research project on artificial intelligence*. Verfügbar unter: <https://raysolomonoff.com/dartmouth/boxa/dart564props.pdf>
- Olbrich, A. (2008). *ITIL kompakt und verständlich. Effizientes IT Service Management - den Standard für IT-Prozesse kennenlernen, verstehen und erfolgreich in der Praxis umsetzen* (Studium, 4., erweiterte und verbesserte Auflage). Wiesbaden: Vieweg+Teubner.
- Österle, H., Becker, J., Frank, U., Hess, T., Karagiannis, D., Krcmar, H. et al. (2010). Memorandum zur gestaltungsorientierten Wirtschaftsinformatik. *Schmalenbachs Zeitschrift für betriebswirtschaftliche Forschung*, 62(6), 664–672. <https://doi.org/10.1007/BF03372838>
- Oronowicz, J., Ley, C., Pachowsky, M., Seil, R., Tischer, T. (2023). Möglichkeiten und Perspektiven zum Einsatz der künstlichen Intelligenz in der Sportorthopädie, Sports Orthopaedics and Traumatology (Volume 39, Issue 1) <https://doi.org/10.1016/j.orthtr.2022.12.002>.
- Peffers, K., Tuunanen, T., Rothenberger, M. A. & Chatterjee, S. (2007). A Design Science Research Methodology for Information Systems Research. *Journal of Management Information Systems*, 24(3), 45–77. <https://doi.org/10.2753/MIS0742-1222240302>
- Powell, J. (2019). Trust Me, I'm a Chatbot: How Artificial Intelligence in Health Care Fails the Turing Test. *Journal of Medical Internet Research*, 21(10), e16222. <https://doi.org/10.2196/16222>
- Pröhl, T. & Zarnekow, R. (2019). Die kurze Geschichte des IT-Servicemanagement: Themen und Fragestellungen im Wandel der Zeit. *HMD Praxis der Wirtschaftsinformatik*, 56(2), 277–288. <https://doi.org/10.1365/s40702-019-00497-7>
- Randus IT, C. (2022). *ITIL 4 Foundation*. Version 1.11, basierend auf ITIL 4.
- Rashid, A. N. M. B., Ahmed, M., Sikos, L. F. & Haskell-Dowland, P. (2020). Cooperative co-evolution for feature selection in Big Data with random feature grouping. *Journal of Big Data*, 7(1). <https://doi.org/10.1186/s40537-020-00381-y>

Rehman, Abdur; An, Hyunbin; Park, Seonghwan; Moon, Inkyu (2024): Automated classification of elliptical cancer cells with stain-free holographic imaging and self-supervised learning. In: *Optics & Laser Technology* 174, S. 110646. DOI: 10.1016/j.optlastec .2024.110646.

Rosati, F. (2023). *The state of the art of AI: challenges and future prospects*. Verfügbar unter: <https://makerfairerome.eu/en/the-state-of-the-art-of-ai-advancements-challenges-and-future-prospects/#>

Styria IT. (2023, 24. Oktober). *Organigramm der Styria IT*. Verfügbar unter: <https://mystyria.net/de/organisationsstruktur-media-it-3475>

Suyal, M. & Goyal, P. (2022). A Review on Analysis of K-Nearest Neighbor Classification Machine Learning Algorithms based on Supervised Learning. *International Journal of Engineering Trends and Technology*, 70(7), 43–48. <https://doi.org/10.14445/22315381/IJETT-V70I7P205>

Uddin, S., Haque, I., Lu, H., Moni, M. A. & Gide, E. (2022). Comparative performance analysis of K-nearest neighbour (KNN) algorithm and its different variants for disease prediction. *Scientific Reports*, 12(1), 6256. <https://doi.org/10.1038/s41598-022-10358-x>

Wilde, T. & Hess, T. (2007). Forschungsmethoden der Wirtschaftsinformatik. *WIRTSCHAFTSINFORMATIK*, 49(4), 280–287. <https://doi.org/10.1007/s11576-007-0064-z>

Winter, R. [Robert] & Baskerville, R. (2010). Science of Business & Information Systems Engineering. *Business & Information Systems Engineering*, 2(5), 269–270. <https://doi.org/10.1007/s12599-010-0123-7>

Wuttke, L. (2023). *Anwendungsgebiete von künstlicher Intelligenz*. Verfügbar unter: (<https://datasolut.com/anwendungsgebiete-von-kuenstlicher-intelligenz/>)

Ye, Y., Zhang, Y. & Zhu, Y. (2022). Exploring the form of big data products and the supporting systems. *Journal of Big Data*, 9(1). <https://doi.org/10.1186/s40537-022-00604-4>