

Datenmanagement im Finanz- und Rechnungswesen



Geleitwort



Daten sind längst zu einem zentralen Erfolgsfaktor für Unternehmen geworden.

Gerade im Finanz- und Rechnungswesen zeigt sich dies besonders deutlich: Verlässliche Daten bilden die Grundlage für korrekte Abschlüsse, aussagekräftige Reports, fundierte Führungsentscheidungen und eine wirksame Unternehmenssteuerung. Gleichzeitig nehmen Datenvolumen, Systemvielfalt und regulatorische Anforderungen laufend zu. Begriffe wie Data Analytics, Business Intelligence, Data Lake, Robotic Process Automation oder Künstliche Intelligenz sind deshalb keine Zukunftsthemen mehr, sondern prägen bereits heute die tägliche Arbeit im Rechnungswesen, im Controlling und in der Treuhandpraxis.

Mit der vorliegenden Broschüre **«Datenmanagement im Finanz- und Rechnungswesen»** möchte SwissAccounting Orientierung bieten. Sie soll Grundlagen vermitteln, Zusammenhänge aufzeigen und vor allem den praktischen Nutzen eines professionellen Datenmanagements sichtbar machen. Denn Digitalisierung beginnt nicht erst bei Dashboards, Automatisierung oder KI-Modellen. Sie beginnt bei sauberen, modellierten, nachvollziehbaren und sicheren Daten. Wer Datenqualität, Datenarchitektur, Datenschutz und Datensicherheit nicht im Griff hat, wird auch die Chancen neuer Technologien nur begrenzt nutzen können.

Als grösster Schweizer Berufs- und Fachverband im Finanz- und Rechnungswesen sieht sich SwissAccounting in der Verantwortung, seine Mitglieder bei diesen Entwicklungen aktiv zu begleiten. Unsere Berufsgruppe steht an einer wichtigen Schnittstelle: Sie verarbeitet, beurteilt und interpretiert finanzielle Informationen und schafft damit Vertrauen – innerhalb der Unternehmen, gegenüber Eigentümern, Behörden, Kapitalgebern und weiteren Anspruchsgruppen. Gerade deshalb wollen wir die digitale Transformation verstehen, einordnen und aktiv mitgestalten.

Die Broschüre ist bewusst praxisnah angelegt. Sie richtet sich an Fachleute im Finanz- und Rechnungswesen, im Controlling, im Treuhandwesen sowie an Führungskräfte, die sich mit den Grundlagen und

Anwendungsmöglichkeiten des Datenmanagements vertraut machen möchten.

Sie soll nicht jedes technische Detail vertiefen, sondern einen verständlichen Zugang eröffnen und aufzeigen, welche Fragen für die Praxis besonders relevant sind: Welche Daten liegen vor? Wie werden sie klassifiziert? Welche Architektur ist sinnvoll? Wo können Automatisierung und KI unterstützen? Welche Risiken bestehen im Bereich Datenschutz und Datensicherheit? Und welche Rolle kommt der Finanzfunktion in der digitalen Transformation zu?

Ein herzlicher Dank gilt allen Autoren, die mit ihrem Fachwissen, ihrer Praxiserfahrung und ihrem grossen Engagement zu diesem Gemeinschaftswerk beigetragen haben. Besonders danken möchte ich meinem Vorstandskollegen Fabian Meisser, der dieses Projekt mit Leidenschaft geleitet und die Gesamtverantwortung dafür übernommen hat. Eine Publikation dieser Art lebt davon, dass verschiedene Perspektiven zusammenkommen: fachliche, technische, rechtliche und organisatorische. Nur so entsteht ein Gesamtbild, das der Breite und Dynamik des Themas gerecht wird.

Die Broschüre erscheint bewusst nicht in gedruckter Form, sondern als Download auf unserer Webseite. Auch dies ist Ausdruck des Themas selbst: Datenmanagement, Digitalisierung und Künstliche Intelligenz entwickeln sich schnell weiter. Eine digitale Publikation erlaubt es uns, Inhalte bei Bedarf zu aktualisieren, neue Entwicklungen aufzunehmen und damit näher an der Praxis zu bleiben.

Ich wünsche dieser Broschüre eine weite Verbreitung. Möge sie dazu beitragen, das Bewusstsein für die Bedeutung eines professionellen Datenmanagements im Finanz- und Rechnungswesen zu stärken, Diskussionen anzuregen und konkrete Impulse für die tägliche Arbeit zu geben.

Prof. em. Dr. Dieter Pfaff
Präsident SwissAccounting

Autoren (in der Reihenfolge der Kapitel)



Fabian Meisser

Fabian Meisser ist Controller mit über 10 Jahren Praxis und ausgebildeter Data Scientist. Als Mitgründer und Geschäftsführender Partner der DataVision AG begleitet er Unternehmen bei der Digitalisierung des Finanz- und Rechnungswesens und entwickelt praxisnahe, datengetriebene Lösungen.



Irfan Yasar

Irfan Yasar ist Head of Controlling, Data Solutions & BI bei der Schweizerischen Post und Gründer der Reporting Factory GmbH. Als eidg. dipl. Experte in Rechnungslegung und Controlling doziert er bei Swiss-Accounting und leitet die Fachkommission Datenmanagement.



Ralph Hutter

Ralph Hutter, eidg. dipl. Informatiker und EMBA, ist Studiengangsleiter an der HWZ Zürich und Head Ecosystems & Partnerships bei Finnova. Er verbindet über 20 Jahre Praxis in Digital Product Management mit akademischer Lehre zu digitalen Geschäftsmodellen, Plattformen und AI.



Jon Cajacob

Jon Cajacob ist Partner bei der DataVision AG und hat zahlreiche BI- & Analytics-Projekte im Finanzbereich erfolgreich umgesetzt. Mit einem Hintergrund in Corporate Finance und fundiertem Datenwissen realisiert er branchenübergreifend Projekte bis hin zu Gesamtarchitekturen auf modernen Datenplattformen.



Abbas Tutcuoglu

Abbas Tutcuoglu beschäftigt sich seit vielen Jahren mit der praktischen Anwendung von KI. Heute arbeitet er im Technical Success Team von OpenAI. Zuvor sammelte er Erfahrung in KI-Strategie und -Transformation bei McKinsey/QuantumBlack sowie als Data Scientist bei Eraneos.



Markus Speck

Markus Speck, dipl. Experte in Rechnungslegung und Controlling. 25 Jahre internationale Fach- und Führungserfahrung in einem Industrie-Konzern. Seit 2010 selbständiger Unternehmensberater mit Schwerpunkt finanzielle Führungssysteme und dem dazugehörigem Prozess- und Daten-Management. Dozent für höhere Fachprüfungen und Lehrmittel-Autor.

Inhalt

Einleitung	5	5 Künstliche Intelligenz	57
		5.1 Was ist künstliche Intelligenz?	57
		5.2 Die verschiedenen Formen der KI	59
		5.3 Entstehung von KI-Modellen	61
		5.4 LLMs	66
		5.5 Anfälligkeiten eines KI-Modells	72
		5.6 Buy-vs.-Build	74
1 Grundlagen: Daten- und Datenbanken	7	6 Datensicherheit und Datenschutz	78
1.1 Daten und Datenbanken	7	6.1 Strategische Grundlagen: Daten als Finanzwert und Risiko	78
1.2 Klassifizierung von Daten	9	6.2 Datenschutz	79
1.3 Das ERP	12	6.3 Datensicherheit	84
1.4 Die Cloud	14	6.4 Cyber Security	86
		6.5 Incident Management	90
		6.6 Business Continuity Management (BCM)	92
		6.7 Der CFO-Blickwinkel	93
2 Data Analytics und Business Intelligence	17	7 Digitale Transformation und Organisation	95
2.1 Definitionen und Reifegrade	17	7.1 Ambition und Zielbild	95
2.2 Bedeutung von Business Intelligence	19	7.2 Führungsverantwortung und Data-Governance	96
2.3 Das BI-Framework	19	7.3 System- und Datenlandschaft	98
2.4 Formen des finalen Reportings	24	7.4 Projektmanagement und agile Methoden	100
2.5 Toolübersicht	25	7.5 Change-Management als zwingender Erfolgsfaktor	103
		7.6 Zusammenfassung	105
3 Datenarchitektur	27	Glossar	106
3.1 Überblick	27	Literatur- und Quellenverzeichnis	110
3.2 Data Warehouse, Data Lake, Lakehouse	29		
3.3 Die hohe Relevanz von Self-Service und Low-Code Tools	33		
3.4 Umsetzung einer modernen Datenarchitektur mit einer Plattformlösung	36		
3.5 Datenqualität: Strategien zur kontinuierlichen Verbesserung	40		
4 Weitere Automatisierungsformen und Tools	42		
4.1 Prozesse und Automatisierungspotenzial	42		
4.2 Robotic Process Automation (RPA)	44		
4.3 Cloud- und Toolbasierte Automatisierung	50		
4.4 Praktische Anwendungsfälle für KMU	51		
4.5 Herausforderungen und Erfolgsfaktoren	54		

Einleitung



Fabian Meisser

Die klassische Aus- und Weiterbildung im Finanz- und Rechnungswesen konzentriert sich weitestgehend auf fachliche Aspekte. Schauen wir uns ein Beispiel an, wie es Ihnen vielleicht in einem Controlling-Aufgabenbuch begegnen könnte:

Aufgabe 1:

Gegeben sind folgende Umsätze und Deckungsbeiträge von Kunden Ihres Unternehmens. Berechnen Sie anhand der im Theoriebuch erlernten Methodik die DB1-Marge in % und nehmen Sie ein Kunden-Clustering vor:

Kunde	Umsatz (CHF)	Direkte Herstellkosten (CHF)
Lehner Möbel AG	250 000	75 000
Alpha Electronics AG	180 000	54 000
Gamma Textil GmbH	320 000	96 000
Steiner Bau AG	410 000	123 000
Helios Pharma SA	275 000	82 500
Nova Foods GmbH	360 000	108 000
Zürich IT Solutions	295 000	88 500
Alpenhotel AG	220 000	66 000
SwissClean Services	190 000	57 000
Bergsport Meier AG	330 000	99 000

Abbildung 1: Typische Ausgangslage einer Controlling-Aufgabe

Konzeptionelle Grundlagen sind wichtig – aber in der Praxis greifen sie zu kurz. Niemand, schon gar nicht Ihr Auftraggeber, wie zum Beispiel der CEO¹, legt Ihnen die fertigen Zahlen auf den Tisch. Ihre Aufgabe beginnt viel früher: Sie müssen die Daten, die Sie in solchen Aufgaben als Ausgangslage bekommen, selbst beschaffen – effizient, schnell und ohne jedes Mal bei Null zu starten. Und selbst wenn Sie die Auswertung erstellt haben, ist sie selten endgültig: Die Auftraggeberin möchte vielleicht zusätzliche Details, eine

andere Periodensicht oder weitere, verwandte Kennzahlen. Das zeigt, wie entscheidend heute eine neue Kompetenz ist: der souveräne Umgang mit Unternehmensdaten – oder neudeutsch: Data Literacy.

Data Literacy geht weit über ein Informatik-Zusatzfach hinaus, in dem man die Grundlagen von Bits und Bytes lernt. Hilfreich ist dieses Wissen zwar, aber die wahre Herausforderung liegt in der Verknüpfung von Finanzwissen und Datenkompetenz, um in der Praxis Ergebnisse zu liefern. Die Datenflut wächst durch immer mehr Tools und Online-Services. Das macht die Aufgabe komplexer. Diese Broschüre richtet sich an Fach- und Führungspersonen im Finanz- und Rechnungswesen, die täglich mit Daten arbeiten – oft in Excel – und nun den nächsten Schritt gehen wollen.

Im Folgenden zeigt diese Broschüre, wie Sie diese Datenkompetenz systematisch aufbauen können. Die einzelnen Kapitel orientieren sich dabei am typischen Reifeweg im Finanz- und Rechnungswesen – von den Grundlagen bis zur organisatorischen Verankerung.

Gibt es eine Methode, mit der Daten auf Knopfdruck aktuell und konsistent zusammengeführt werden können – auch wenn sie aus verschiedenen Systemen stammen? Eine Lösung dafür kann Business Intelligence (BI) sein. Nach einer Einführung in die Grundbegriffe der Datenwelt zeigt die Broschüre, wie BI Ihre Auswertungsarbeit automatisieren kann. Spätestens wenn Ihr Unternehmen jedoch mehrere Finanzabteilungen pro Division, Ländergruppe oder Produktbereich hat, braucht es Konzepte, die die zusätzliche Komplexität abdecken können. Dann kommen Methoden wie Data Warehouse und Data Lake ins Spiel. Der Autor dieser Einleitung wurde kürzlich gefragt: «Für mich als CEO ist so ein Data Warehouse oder Data Lake dasselbe, oder muss ich da mehr wissen?» Die Antwort lautet: vielleicht. Sicher ist jedoch: Für Sie als Fachperson im Finanz- und Rechnungswesen, die

¹Aus Gründen der besseren Lesbarkeit wird in dieser Broschüre auf geschlechterspezifische Doppelnennungen und Sonderformen (z. B. Anwender*innen) verzichtet. Sämtliche Personenbezeichnungen gelten gleichermassen für alle Geschlechter.

Einleitung

finanzielle Führung in den 2030er-Jahren und darüber hinaus betreiben werden, ist es entscheidend, diese Konzepte auseinanderhalten zu können.

Neben automatisierten Reportings gibt es weitere Automatisierungen, die für Fachleute relevant sind: Fakturaversand, Stammdateneingaben, Einlesen von Buchungsjournalen und viele mehr. Dem Automatisierungsthema wird ein eigenes Kapitel gewidmet – ebenso der künstlichen Intelligenz. Wir erläutern die Grundlagen und zeigen konkrete Anwendungsfälle im Accounting. Doch technische Möglichkeiten sind wertlos, wenn Sie diese wegen Sicherheitsbedenken nicht nutzen dürfen. Sie brauchen das richtige Vokabular, um mit IT-Security-Verantwortlichen zu sprechen, und müssen Ihr Team für Risiken sensibilisieren.

Zum Schluss stellt sich die Frage: Wie organisieren Sie all das? Wo beginnen Sie? Wie nehmen Sie Ihr Team mit? Die abschliessenden Ausführungen zu diesen Fragen finden Sie im Kapitel zur digitalen Transformation und Organisation.

Grundlagen: Daten- und Datenbanken

Fabian Meisser



1.1 Daten und Datenbanken

Gemäss Wikipedia sind Daten Zeichen, die eine Information darstellen, und Datenbanken sind eine Sammlung von Daten. Datenbanken gibt es nach dieser Definition in der Menschheitsgeschichte schon seit Jahrhunderten. Auch Bücher und Dictionaries repräsentieren Datensätze in physischer Form. In der modernen Arbeitswelt assoziieren wir Datenbanken häufig mit auf Computern gespeicherten und schnell zugänglichen Daten. Bereits in den 1970er Jahren veröffentlichte Oracle eine relationale Datenbank; der heute immer noch breit eingesetzte SQL-Server hat seinen Ursprung im Jahr 1989.² Eine wichtige Entwicklung – und bis heute ein zentraler Begriff – ist die relationale Datenbank.³ Sie bezeichnet einen Datenbanktyp, der Daten in Zeilen und Spalten organisiert, die zusammen Tabellen bilden, welche über Beziehungen miteinander verknüpft sind. Um dies zu konkretisieren, zeigt Abbildung 2 das Beispiel eines Kassensystems (Ausschnitt) eines Gastronomieunternehmens:

RECNUM	DATE	PRODID	PRICE	Amount	PAYAMOUNT	PAYAMOUNT_EX_VAT
102334	09.12.2025	321	32.25	1	32.25	29.83
102335	09.12.2025	402	4.75	1	4.75	4.39
102336	09.12.2025	402	4.75	1	4.75	4.39
102337	09.12.2025	321	32.25	1	32.25	29.83
102338	09.12.2025	321	32.25	1	32.25	29.83
102339	14.12.2025	402	4.75	1	4.75	4.39
102340	15.12.2025	321	32.25	1	32.25	29.83
102341	16.12.2025	402	4.75	1	4.75	4.39
102342	17.12.2025	321	25.8	1	25.8	23.87
102341	02.01.2026	402	4.75	1	4.75	4.39

Abbildung 2: Transaktionstabelle aus einem Kassensystem⁴

PRODID	Name	LISTPRICE
321	Sushi Plate Medium	32.25
402	O-Saft Frisch 3 DL	4.75
501	Galce Vanille 1 Port	6.5

Abbildung 3: Ausschnitt aus einer Produkt-Stammdatentabelle

Wichtige Begriffe hierbei:

- **Tabellen** bestehen aus **Zeilen (Datensätze)** und **Spalten (Attribute, in der Praxis auch «Felder»** genannt).
- Jede Zeile repräsentiert einen einzelnen Datensatz, jede Spalte ein Merkmal dieses Datensatzes.
- Tabellen sind über **Schlüssel** miteinander verknüpft:
 - **Primärschlüssel:** eindeutige Kennung für jeden Datensatz. (Spalte «PRODID» in Tabelle 3)
 - **Fremdschlüssel:** verweist auf den Primärschlüssel einer anderen Tabelle, um Beziehungen herzustellen. (Spalte «PRODID» in Tabelle 2)

Wie müssen wir uns dies in der physischen Welt vorstellen – wo sind diese Daten gespeichert? Das klassische Modell ist das Server Client Computing: Datenbanken sind in der Praxis oft gross und würden die Kapazität eines Personal Computers sprengen und ein paralleles Arbeiten verunmöglichen. Daher wird die Datenbank auf einem Server gespeichert (man sagt auch «gehostet»). Letztlich ist ein Server ebenfalls ein Computer mit RAM, Prozessoren und allem, was dazugehört; er hat jedoch nicht den Zweck, morgens von einer Mitarbeiterin oder einem Mitarbeiter gestartet und nur für den persönlichen Gebrauch genutzt zu werden, sondern zentral die Datenbank und ihre Funktionen zur Verfügung zu stellen (Server «dienen», Clients lassen sich bedienen).

Wenn Sie sich morgens in Ihr SAP, Abacus, Dynamics oder ein anderes «Lieblingssystem» einloggen, sitzen Sie auf der Client-Seite, und der Server «dient» Ihnen: Transaktionen, die Sie eingeben, werden auf dem Server und der darauf installierten Datenbank gespeichert und können von Ihnen oder Ihren Kolleginnen und Kollegen (mit ihren eigenen Clients) wieder auf diese Transaktionen zugreifen.

² Harrison, 2015, S. 4

³ <https://www.ibm.com/think/topics/relational-databases>

⁴ Eigene Darstellung, erfundene Beispieldatensätze

1 Grundlagen: Daten- und Datenbanken

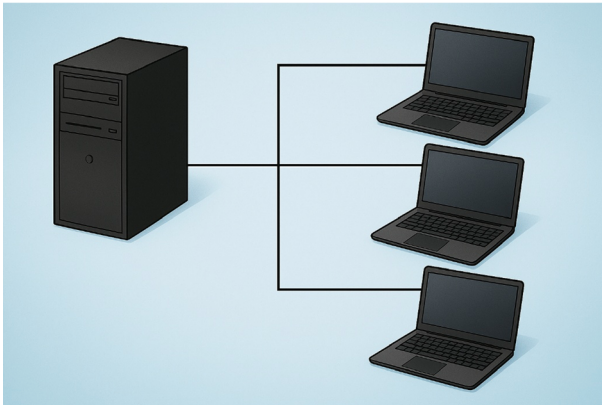


Abbildung 4: Server (links), der über ein Netzwerk mit 3 Clients verbunden ist⁵

Typische Datenbanken, mit denen Mitarbeitende im Accounting zu tun haben, sind:

- ERP (Enterprise Resource Planning)
- CRM (Customer Relationship Management)
- EPM (Enterprise Performance Management)
- Vertragsdatenbanken
- Treasury-Systeme
- Investmentplattformen

Wie interagiert jetzt der Finanzmitarbeitende mit diesen Daten? Nehmen wir das Beispiel unserer relationalen Finanzdatenbank oben.

Uns interessiert die Summe des Umsatzes (Attribut «PAYAMOUNT_EX_VAT»); hierzu kann SQL (Structured Query Language) als Abfragesprache verwendet werden:

```
SELECT
    SUM(PAYAMOUNT_EX_VAT) AS Umsatz
FROM Transaktionen;
```

Resultat Umsatz = 165.14

Jetzt möchten wir noch einen sprechenden Namen dazu und nur den Dezember 2025 auswerten. Dazu muss über den Schlüssel der Name aus der separierten Stammdatentabelle der Konten extrahiert und auf den Dezember 2025 eingeschränkt werden:

```
SELECT
    p.Name,
    SUM(t.PAYAMOUNT_EX_VAT) AS Umsatz
FROM Transaktionen AS t
JOIN Products AS p
    ON p.PRODID = t.PRODID
WHERE t.DATE >= DATE '2025-12-01'
    AND t.DATE < DATE '2026-01-01'
GROUP BY p.Name
ORDER BY p.Name;
```

Resultat:

Name	Umsatz
0-Saft Frisch 3 DL	17.56
Sushi Plate Medium	143.19

Dies war nun die technische Betrachtung. Als Fachperson Finanzen arbeiten Sie mit einer Applikation, in der Sie zum Beispiel über eine Maske einen Kontoauszug erstellen. Im Hintergrund läuft eine Datenbankabfrage wie oben – nur benutzerfreundlich aufbereitet. Das, was der Endbenutzer sieht, nennt man «UI» (User Interface); begrifflich zu unterscheiden vom «Backend», also dem, was im Hintergrund in der Datenbank inklusive der damit verbundenen Prozesse passiert.

⁵Eigene Darstellung, Quelle: Microsoft Copilot

1 Grundlagen: Daten- und Datenbanken

1.2 Klassifizierung von Daten

Im alltäglichen Umgang mit Daten ist es hilfreich, sich bewusst zu machen, mit welcher Art von Daten wir es in einem bestimmten Kontext gerade zu tun haben. Dabei hilft, dass Daten verschieden klassifiziert werden können. Nachfolgend vier der gängigsten Klassifizierungsarten von Daten.

Klassifizierung	Ausprägungen	Beispiele
Struktur	Strukturierte Daten	Datenbanktabellen
	Semi-strukturierte Daten	JSON
	Unstrukturierte Daten	Wordtextdatei
Entstehung und Nutzung	Stammdaten	Kundenstammdaten
	Transaktionsdaten	Verkaufstransaktionen
	Metadaten	Erstellungsdatum
Zeitbezug	Echtzeitdaten	Maschinensensordaten
	Batchdaten	Buchhaltungsdaten
Datentypen	Integer (Ganzzahl)	5
	Dezimal	5.5
	Text	Texte
	Datum	01.10.2027

Abbildung 5: Klassifizierung von Daten (eigene Darstellung)

Klassifizierung nach Struktur

Strukturierte Daten sind strikt geordnet und folgen einem Schema mit eindeutigen Datentypen. Eine Datenbank erzwingt durch ihr hinterlegtes Schema, dass Daten strukturiert vorliegen. Auch in Excel können strukturierte Daten vorliegen, z.B. wenn eine Tabelle als Tabelle formatiert wird (Einfügen → Tabelle).

	A	B	C	D	E	F
1	SalesKey	Date	ProductKey	UnitPrice	SalesQuantity	Turnover
2	1255711	01.01.2020	348.00	385.97	43	16'596.71
3	416607	01.01.2020	787.00	7.78	41	318.98
4	486052	01.01.2020	1'586.00	7.66	38	291.08
5	152343	01.01.2020	561.00	76.06	5	380.30
6	946470	01.01.2020	718.00	88.87	22	1'955.14
7	364196	01.01.2020	489.00	574.24	98	56'275.52
8	266739	01.01.2020	655.00	109.06	23	2'508.38
9	665197	01.01.2020	620.00	129.50	25	3'237.50
10	1022292	01.01.2020	552.00	1'656.84	2	3'313.68
11	499217	01.01.2020	712.00	91.91	41	3'768.31
12	433264	01.01.2020	1'044.00	446.74	39	17'422.86
13	1140	01.01.2020	1'297.00	18.27	38	694.26
14	443256	01.01.2020	785.00	8.46	100	846.00
15	929145	01.01.2020	405.00	591.56	28	16'563.68
16	153532	01.01.2020	1'093.00	333.13	53	17'655.89
17	167044	01.01.2020	1'372.00	21.11	74	1'562.14
18	2435	01.01.2020	1'433.00	158.90	4	635.60
19	1024149	01.01.2020	710.00	72.43	92	6'663.56
20	1036697	01.01.2020	1'323.00	20.52	84	1'723.68
21	13503	01.01.2020	568.00	356.98	81	28'915.38
22	1221393	01.01.2020	1'075.00	249.85	59	14'741.15
23	28918	01.01.2020	1'185.00	650.50	77	50'088.50
24	522255	01.01.2020	953.00	118.44	7	829.08
25	1169028	01.01.2020	1'484.00	131.04	35	4'586.40
26	544723	01.01.2020	52.00	161.92	8	1'295.36
27	1009200	01.01.2020	601.00	428.63	51	21'860.13
28	340358	01.01.2020	377.00	367.31	9	3'305.79
29	1200519	01.01.2020	1'421.00	180.96	87	15'743.52
30	816505	01.01.2020	658.00	153.55	88	13'512.40

Abbildung 6: Strukturierte Daten in einem Excel-Spreadsheet

Aus Sicht einer Datenbank würde dieses Excel noch nicht alle Strukturanforderungen erfüllen, um wirklich «strukturiert» im engeren Sinne zu sein. So müssten beispielsweise Datentypen definiert werden (siehe weiter unten), und es wären Spalten vorgesehen, die im aktuellen Anwendungsfall zwar leer bleiben, später jedoch für andere Zwecke benötigt werden könnten.

Die meisten Excel-Files in der Praxis sind jedoch semi-strukturiert. Datenhaltung und Auswertung sind oft nicht klar getrennt; es gibt Freitext-Titel, Zwischentotale, Abstandszeilen und Bemerkungen innerhalb von Zellen. Immerhin befindet sich jeder Wert in einer klar bestimmten Zelle, z.B. E7 – es gibt also eine gewisse Organisation der Daten, sie folgen jedoch keinem klaren Schema.

	A	B	C	D	E	F	G
1							
2							
3		Umsatz nach Produkt per 30.9.2026					
4		Verkaumsatz in Berichtswahrung CHe					
5		Produkt	Asia	Europe	North America	Total	
6		Adventure Works Laptop16 M1601 Blue	119'638	280'605	1'837'774	2'238'017	
7		Adventure Works Laptop16 M1601 Silver	388'731	492'744	1'909'923	2'789'398	
8		Zwischentotal Laptop	508'370	773'349	3'747'697	5'029'416	
9		Contoso Desktop Alternative Bundle E200 Grey	610	5	10'884	11'499	
10		Contoso Desktop Alternative Bundle E200 White	1'850	1'763	11'279	14'892	
11		Fabrikam SLR Camera M149 Pink	21'662	63'698	376'273	461'634	
12		The Phone Company Touch Screen Phones SAW/On-wall M806 Black	166'144	128'439	931'331	1'226'914	
13		WWI 2GB Pulse Smart pen M100 Black	97'084	137'787	519'808	754'679	
14		Gesamtergebnis	793'719	1'078'042	5'991'472	7'863'234	
15		in % vom Vorjahr	-9.09%	7.53%	9.89%	7.17%	
16		Hinweis: ohne Produkt nach Dekonsolidierung per 1.7.2026					
17							

Abbildung 7: Semi-Strukturierte Daten in einem Excel-Spreadsheet

1 Grundlagen: Daten- und Datenbanken

Ein klassisches Beispiel für semi-strukturierte Daten sind Übertragungsformate wie JSON oder XML. Diese haben – anders als typischerweise in einer Datenbank – kein fixes, universelles Schema; sie benötigen zwar ein Schema, dieses kann jedoch je nach Anwendungsfall individuell sein, innerhalb gewisser Notationen und Grundregeln. Diese Schemas sind im Internetzeitalter immer relevanter geworden, weil es Standards für die Übertragung von Daten über das Internet und für Schnittstellen braucht. Wenn also Ihr Enterprise Resource Planning (ERP) Buchungssätze oder Stammdaten im XML-Format aufnehmen kann und Ihr Spesentool XML ausgeben kann, haben die beiden Systeme eine gemeinsame Kommunikationsmöglichkeit, auch wenn die Datenbankschemata der Systeme selbst deutlich unterschiedlich sind. Nachfolgend ein Beispiel für ein JSON, das Spesendaten ausgibt

```
{
  "spesen": [
    {
      "datum": "2026-01-17",
      "mitarbeiternummer": "083",
      "kostenstelle": "4102",
      "betrag": 23.40
    },
    {
      "datum": "2026-06-03",
      "mitarbeiternummer": "214",
      "kostenstelle": "7320",
      "betrag": 37.95
    },
    {
      "datum": "2026-11-22",
      "mitarbeiternummer": "007",
      "kostenstelle": "1057",
      "betrag": 49.10
    }
  ]
}
```

In einer Datenbank wären jeweils alle Datenfelder vorhanden; Inhalte, die in einem bestimmten Kontext nicht relevant sind, bleiben einfach leer. Das oben gezeigte JSON-Schema hingegen definiert lediglich die vier für die Übertragung relevanten Felder (Datum, Mitarbeiternummer, Kostenstelle, Betrag).

Jetzt dieselben Daten im XML-Format:

```
<spesen>
  <datensatz>
    <datum>2026-01-17</datum>
    <mitarbeiternummer>083</mitarbeiternummer>
    <kostenstelle>4102</kostenstelle>
    <betrag>23.40</betrag>
  </datensatz>
  <datensatz>
    <datum>2026-06-03</datum>
    <mitarbeiternummer>214</mitarbeiternummer>
    <kostenstelle>7320</kostenstelle>
    <betrag>37.95</betrag>
  </datensatz>
  <datensatz>
    <datum>2026-11-22</datum>
    <mitarbeiternummer>007</mitarbeiternummer>
    <kostenstelle>1057</kostenstelle>
    <betrag>49.10</betrag>
  </datensatz>
</spesen>
```

Und als letztes Beispiel das deutlich einfacher zu lesende CSV-Format:

```
datum,mitarbeiternummer,kostenstelle,betrag
2026-01-17,083,4102,23.40
2026-06-03,214,7320,37.95
2026-11-22,007,1057,49.10
```

Es besteht kein Anlass, vor technisch anmutenden Daten Ehrfurcht zu haben – letztlich handelt es sich einfach um andere Definitionen – vergleichbar mit Sprachen, die jeweils eigenen Regeln folgen.

Als letzte Kategorie bleiben die unstrukturierten Daten. Beispiele finden sich unmittelbar: Die vorliegende PDF-Broschüre enthält einen Mix aus Texten, Tabellen und Grafiken und lässt sich nicht ohne Weiteres in Datenbankfelder einlesen oder in ein JSON-Format übertragen. Typische unstrukturierte Daten sind zudem Bild-, Video- und Audio-dateien. Um die Beispielreihe in Excel abzuschließen, betrachten wir ein selbsterklärendes Beispiel für unstrukturierte Daten in Excel.

1 Grundlagen: Daten- und Datenbanken

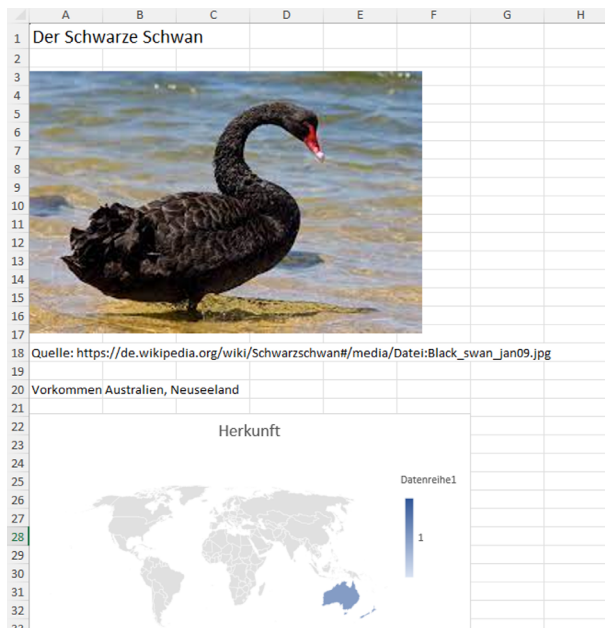


Abbildung 8: Unstrukturierte Daten in einem Excel-Spreadsheet

Klassifizierung nach Entstehung

Daten können neben einer technischen Betrachtung auch inhaltlich klassifiziert werden. Für die später folgenden Kapitel zu Business Intelligence und Data Warehouse ist die Unterscheidung zwischen **Transaktionsdaten** und **Stammdaten** zentral.

Transaktionsdaten sind beispielsweise der General Ledger eines Buchhaltungssystems, in dem jede Buchung aufgeführt ist. Dort sehen wir häufig nur Nummern (z. B. die Kontonummer «4650»), jedoch keine Namen. Diese – und viele weitere beschreibende Elemente – finden sich in den Stammdaten.

Damit veranschaulichen wir zugleich die oben eingeführten Begriffe Primärschlüssel und Fremdschlüssel (Konto Nr.).

Metadaten als weitere Kategorie sind Datensätze, die typischerweise nicht aktiv eingegeben werden, sondern im Prozess der Datenerhebung mitgeführt werden. Wenn eine Buchhalterin oder ein Buchhalter eine Buchung im System erfasst, werden für die Buchung im folgenden Beispiel fünf Werte eingegeben: Datum, Konto im Soll, Konto im Haben, Betrag und Text.

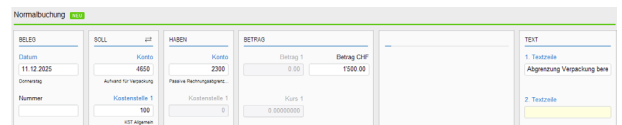


Abbildung 9: Eingabemaske eines Buchhaltungssystems

Im Hintergrund werden gleichzeitig automatisch weitere Metadaten erfasst und in die Datenbank geschrieben:

- Buchungserfasser
- Buchungszeitstempel
- ERP-Versionsnummer
- Technische Transaktionsnummer
- Herkunftsquelle (manuelle Buchung)
- etc.

Klassifikation nach Zeitbezug

Beim Umgang mit Daten begegnen Ihnen oft Begriffe wie **«Realtime»** oder Echtzeitdaten. Damit ist gemeint, dass Sie nahezu in Echtzeit auf aktuelle Werte zugreifen können. Ein Dashboard zeigt den Produktionsoutput zum Beispiel fortlaufend an. Die Übertragung erfolgt also unmittelbar, technisch bedingt jedoch mit geringer Verzögerung – ähnlich wie bei einem Telefongespräch, in dem die Latenz kaum wahrnehmbar ist. Damit dies funktioniert, müssen die Daten unmittelbar bezogen und an den Ort der Auswertung übertragen werden.

Im Finanzkontext sind hingegen Batch-Daten häufiger anzutreffen. Denken Sie an die Stundenerfassung: Einige Mitarbeitende erfassen während des Tages einzelne Stunden, andere erst am Abend den ganzen Tag, wieder andere am nächsten Tag. Ein Bezug sämtlicher Daten zu einem fixen Zeitpunkt (oft nachts) wird als **Batch**-Daten bzw. «Stapelverarbeitung» bezeichnet.

1 Grundlagen: Daten- und Datenbanken

Klassifikation nach Datentypen

Es gibt – abhängig vom Datenbank System – verschiedene Datentypen. Nachfolgend, exemplarisch ein Auszug der gängigsten:

Datentyp	Beispiel
Integer (Ganzzahl)	125
Dezimal	1.256
Text	Accounting Stars
Date	11.12.2025
Date/time	11.12.2025 23:59:59
binär	0

Abbildung 10: Datentypen

Warum ist dies relevant? Zum einen wird für ein bestimmtes Feld in der Datenbank eine definierte Anzahl Speicher Bytes reserviert. Ein Integer benötigt deutlich weniger Speicher als ein Text – schon deshalb, weil das Alphabet mehr Zeichen umfasst als die Ziffern 0 bis 9, ganz abgesehen von Sonderzeichen. Aus Gründen des Speicherbedarfs und der Performance ist es daher sinnvoll, mit passenden Datentypen zu arbeiten.

Zum anderen sichern Datentypen ein bestimmtes Mass an Datenqualität. In Excel kann eine Zelle grundsätzlich jeden Inhalt aufnehmen – unter Umständen auch ein ungültiges Datum wie den 29.02.2027. Spätestens wenn Datumsdifferenzen berechnet oder Filter nach Jahren und Monaten bereitgestellt werden

sollen, führt dies zu Problemen. Datentypen verhindern solche Inkonsistenzen, indem sie Eingaben auf gültige Formate beschränken.

1.3 Das ERP

Was ist ein ERP? Orientieren wir uns an der Definition, welche eines der weltweit am weitesten verbreiteten ERP-Systeme, SAP, angibt:

«Enterprise Resource Planning (ERP) ist ein Softwaresystem, das Unternehmen bei der Rationalisierung ihrer zentralen Geschäftsprozesse – einschliesslich Finanzen, Personalwesen, Fertigung, Lieferkette, Vertrieb und Beschaffung – mit einer einheitlichen Sicht auf die Aktivitäten unterstützt und eine zentrale Datenquelle bietet.»⁶

Und weiter unten in den FAQ:

«Ein ERP-Softwaresystem besteht aus einer Reihe von integrierten Anwendungen oder Modulen zur Verwaltung der wichtigsten Geschäftsprozesse eines Unternehmens. ERP-Module sind in ein komplettes System integriert und nutzen eine **gemeinsame Datenbank**, um Prozesse und Informationen unternehmensweit zu optimieren. Wenn Unternehmen wachsen, können sie den Umfang ihres ERP-Systems erweitern.»⁷

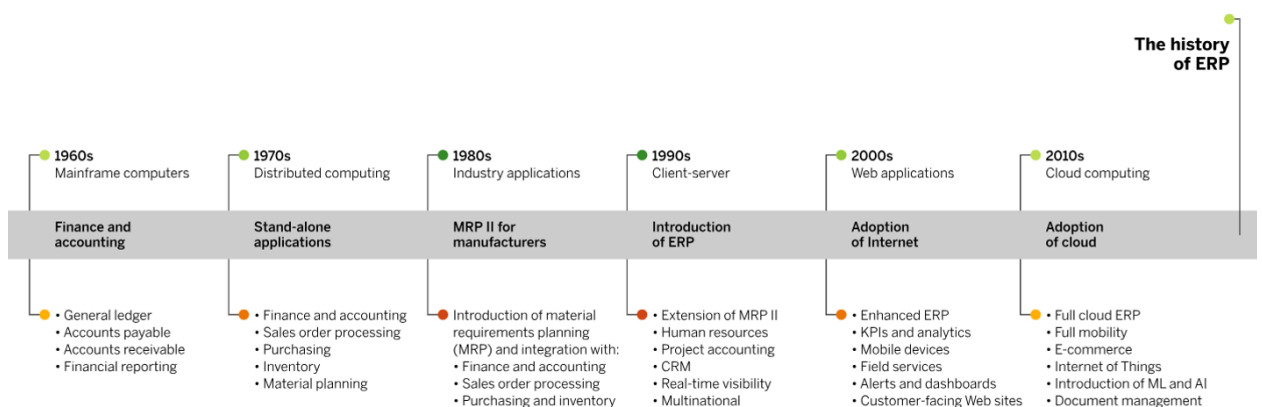


Abbildung 11: Geschichte der ERP-Systeme⁸

⁶ <https://www.sap.com/germany/products/erp/what-is-erp.html>

⁷ <https://www.sap.com/germany/products/erp/what-is-erp.html>

⁸ <https://www.sap.com/germany/products/erp/what-is-erp.html>

1 Grundlagen: Daten- und Datenbanken

Produkteportfolio

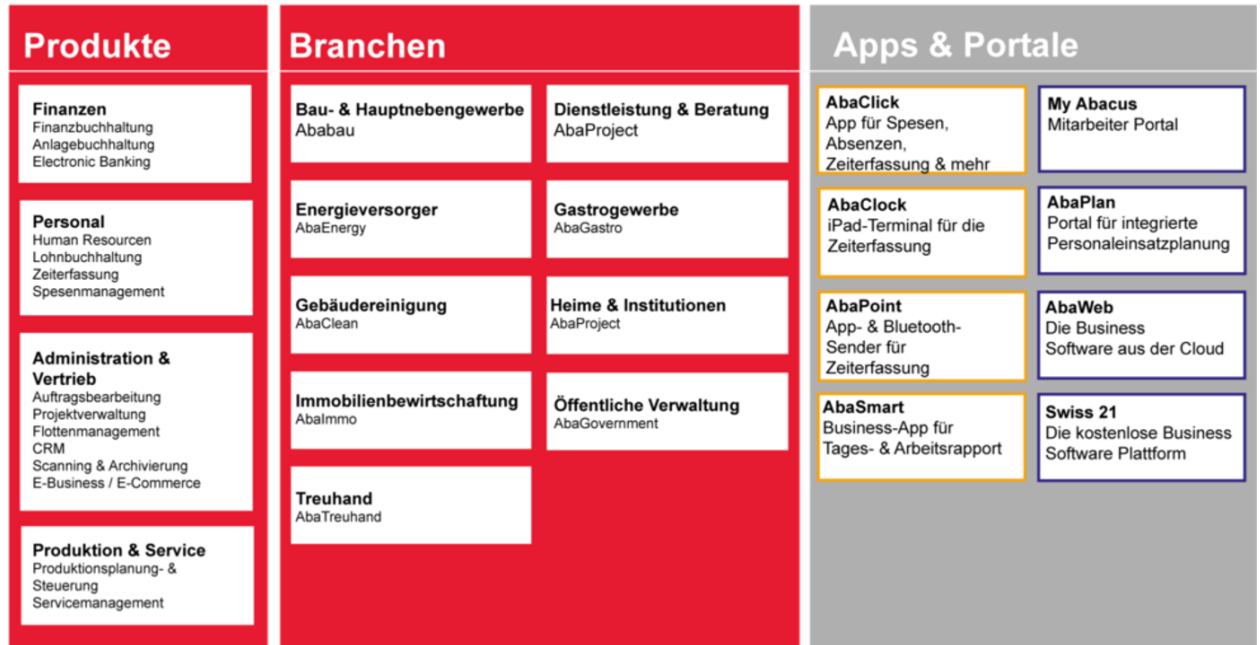


Abbildung 12: Produkteportfolio eines Schweizer ERP-Systems⁹

Kurz gesagt, geht es um die softwaregestützte Abwicklung zentraler Unternehmensprozesse. Der Begriff Datenbank taucht hier ebenfalls auf: Alle Eingaben – seien es Buchungen oder die Pflege von Stammdaten – werden in Datenbanken gespeichert. Eine Finanzmitarbeiterin, die den grössten Teil der Arbeitszeit im ERP verbringt, hat jedoch möglicherweise nie die Datenbanksicht auf das eigene ERP erhalten, weil die Interaktion auf der Benutzeroberfläche, auch User Interface (UI) genannt, und nicht auf der Datenbankebene erfolgt.

Die Geschichte von ERP geht – je nach Betrachtungsweise – bis in die 1960er-Jahre oder weiter zurück. Wie in der folgenden Darstellung ersichtlich, stand am Anfang das Finanzwesen im Zentrum. Der «General Ledger», also die Aufzeichnung des Hauptbuchs einer Buchhaltung, bildet den Kern. Nebenbücher wie Kreditoren oder Debitoren sind typische weitere Bestandteile. Später kamen Material- und Lagerbewirtschaftung, Verkaufsabwicklung, Projektmanagement, Einkauf, CRM (Customer Relationship Management) und vieles mehr hinzu.

Man kann ERP heute also bei Weitem nicht nur mit einer Buchhaltungssoftware gleichsetzen; vielmehr ist die Buchhaltung eines von mehreren Bestandteilen eines ERP, welches zahlreiche weitere Funktionen bereitstellt.

Betrachtet man ein weiteres in der Schweiz typisches ERP – Abacus –, zeigt sich, dass neben den erwähnten Erweiterungen auch verschiedene Branchenlösungen sowie zusätzliche Applikationen und Portale hinzukommen:

Theoretisch könnte man schliessen, dass mit einem ERP für ein Unternehmen alles wunderbar gelöst ist. Für jeden Prozess gibt es ein Modul, das sich einfach beschaffen lässt. Aus Sicht der ERP-Anbieter wäre die ideale Welt, dass ausnahmslos alle Prozesse im selben ERP abgewickelt werden. Die Realität ist jedoch meist anders. Systemlandschaften sind historisch gewachsen, und die Software kann nicht bei jeder Veränderung neu «auf der grünen Wiese» konzipiert werden. Einige Beispiele:

⁹<https://www.bsi-partner.ch/abacus/>

1 Grundlagen: Daten- und Datenbanken

- Unternehmen A ist stark gewachsen und führt ein neues ERP ein. Das vor zwei Jahren eingeführte Salesforce-CRM bleibt bestehen, weil die Kundenprozesse darin nach viel Arbeit optimiert wurden. Obwohl das neue ERP ebenfalls ein CRM bietet, überzeugt dessen Funktionsumfang weniger als die bestehende Lösung.
- Unternehmen B hat SAP im Einsatz und akquiriert das deutlich kleinere Unternehmen C mit Proffix-Buchhaltung. Das Management will die Agilität des neu akquirierten Unternehmens nicht durch einen jahrelangen SAP-Integrationsprozess lähmen. Es entscheidet, für mindestens fünf Jahre zwei ERP-Systeme parallel zu betreiben.
- Für die wachsende Restaurantkette C sind zwei operative Systeme zentral: das Kassensystem vor Ort und das Stunden-/Schichtplanungstool. Man wählt das bestgeeignete Kassensystem. Dieses bildet jedoch weder Buchhaltungsprozesse wie MWST, Lohnbuchhaltung und Kreditoren ab noch hat es Funktionen für die Stundenplanung. Für die Einsatzplanung gibt es ein günstiges Branchen-Tool, das die GAV-Regeln korrekt umsetzt und jährliche Updates im Preis inkludiert. Die Alternative, die Stundenerfassung des neu angeschafften Buchhaltungssystems, bietet dies nicht. Ein teures Customizing wäre nötig und müsste bei Änderungen jährlich nachgezogen werden.
- Unternehmen D hat bisher mit Waren gehandelt. Neu soll auch selbst produziert werden. Das bisherige ERP bietet ein Produktionsmodul, ist aber zehnmal teurer als eine spezialisierte Branchenlösung nur für die Produktion, die zusätzlich fertige Schnittstellen ins vorhandene ERP mitliefert. Die Entscheidung fällt leicht.
- Unternehmen E nutzt seit Jahren dasselbe ERP, hat viel in die Prozessoptimierung investiert und die Mitarbeitenden sind vertraut mit dem System. Der Anbieter ist jedoch weniger innovativ. Kreditorenbelege müssen beispielsweise manuell eingelesen werden, und der Kreditoren-Freigabeprozess gilt als grosse Schwäche. Eine Anbieterin stellt ein KI-gestütztes Kreditorenmodul mit überzeugendem Freigabeprozess vor. Das Unternehmen ersetzt das ERP-integrierte Kreditorenmodul durch die innovative Alternative. Der Rest des ERP bleibt bestehen.

Diese Beispiele zeigen den Zielkonflikt in der Praxis auf: Ein einziges System bringt integrierte Abläufe und reduziert Schnittstellenprobleme. Die Gründe für eine heterogene Systemlandschaft sind jedoch vielfältig und oft überzeugend. Dazu gehören Preisunterschiede, historisch gewachsene Strukturen, Akquisitionen, passgenaue Branchenlösungen, bereits funktionierende Alternativen, die Vermeidung

übermässiger Abhängigkeit und sehr spezifische Unternehmenscharakteristika.

Abschliessend sei erwähnt, dass einige Unternehmen ERP-Funktionen inhouse aufbauen. Kleinere Unternehmen nutzen dafür etwa Access oder die Power Platform zur Prozessabbildung. Es gibt auch Beispiele international tätiger Grossunternehmen, die diesen Weg wählen, etwa Lidl.¹⁰

1.4 Die Cloud

Möglicherweise haben die Personen, welche den Begriff der «Cloud» (Wolke) geprägt haben, der Technologie keinen Gefallen gemacht, denn die Assoziation ist oft: Die Daten sind irgendwo, diffus, nicht greifbar und nicht unter eigener Kontrolle. Diese ersten Assoziationen treffen so nicht zu. Teilweise bestätigt sich einzig der Punkt Kontrolle: Sie ist höher, wenn ein Unternehmen seine Daten auf eigener Infrastruktur betreibt.

Um die Cloud besser zu verstehen, hilft die Frage nach dem Gegenteil. Wenn ein Unternehmen ein System nicht in der Cloud hat, wo dann? Die Alternative ist «On-Premise», also lokal auf eigener Infrastruktur installiert. Die Server, auf denen die relevanten Datenbanken liegen – zum Beispiel die ERP-Datenbank –, stehen im Serverraum des Unternehmens. Die Firma besitzt die Hardware, hält sie instand, führt Updates durch und betreibt die Systeme. Auf dieser Infrastruktur liegen ausschliesslich die eigenen Daten.

Bei der Cloud gehört die Hardware einem externen Anbieter. Dieser stellt den Betrieb sicher, hält die Systeme instand, spielt Updates ein und sorgt für Backups. Die Verbindung zum Unternehmen, das den Cloud-Dienst nutzt, erfolgt über das Internet.

Als Einstieg ins Thema muss man sich also nur zwei Definitionsmerkmale merken: Erstens ist die Hardware ausgelagert. Zweitens erfolgt der Zugriff auf die Daten über das Internet.

¹⁰ <https://www.inside-it.ch/statt-auf-sap-setzt-lidl-auf-eine-individualloesung-20240503>

1 Grundlagen: Daten- und Datenbanken

Es gibt verschiedene Arten von Cloud-Service-Modellen. Typischerweise unterscheidet man drei:

Cloud Model	Beschrieb	Beispiele
Infrastructure as a Service (IaaS)	Anbieter stellt Server, Speicher und Netzwerk bereit. Betriebssystem und Anwendungen werden selbst verwaltet	Microsoft Azure (Virtuelle Maschinen), Amazon EC2, Google Compute Engine
Platform as a Service (PaaS)	Zusätzlich zu IaaS verwaltet der Anbieter auch Betriebssystem, Middleware und Laufzeitumgebung.	Microsoft Azure App Service, Google App Engine
Software as a Service (SaaS)	Zusätzlich zu IaaS und PaaS wird auch die Applikation selbst vom Anbieter betrieben und gewartet.	Microsoft 365, Salesforce, Google Workspace, Dropbox

Abbildung 13: Die drei Cloud-Modelle¹¹

Diese Cloud-Modelle bauen aufeinander auf. Je weiter man in dieser Liste von oben nach unten geht, desto mehr Aufgaben werden an den Cloud-Anbieter ausgelagert.

Zusätzlich ist die Unterscheidung in Private Cloud und Public Cloud wichtig. Bei einer Private Cloud werden die Dienste exklusiv für das Kundenunternehmen bereitgestellt. In einer Public Cloud teilen sich verschiedene Kunden dieselbe physische Infrastruktur, arbeiten aber in getrennten Mandanten mit separaten Logins. Über ein Berechtigungskonzept ist sichergestellt, dass jedes Unternehmen nur die eigenen Daten sieht und bearbeiten kann.

Verknüpfen wir dies mit den einführenden Themen zu Datenbanken und ERP. So könnten KMU in der Schweiz den Betrieb ihres Abacus-Systems gestalten:

- Unternehmen A betreibt einen eigenen Serverraum, in dem es Server und Speicher sowie das Netzwerk selbst managt. Abacus ist lokal installiert. Die Clients greifen intern auf den Server zu, und alle Daten bleiben On-Premise.¹²
- Unternehmen B hat keinen eigenen Serverraum und mietet bei Microsoft Server, physischen Speicher und das Netzwerkmanagement. Windows wird vom Unternehmen selbst installiert und gepflegt, Updates werden eingespielt und Abacus wird eigenständig installiert. Über Remote Desktop besteht voller Zugriff auf die Datenbank und die Applikation. Das entspricht IaaS.
- Unternehmen C verfügt ebenfalls über keinen eigenen Serverraum. Im Unterschied zu B möchte es sich nicht um Betriebssystem, Middleware und weitere technische Komponenten kümmern und bezieht diese als Dienst mit. Die eigene Abacus-Installation läuft auf dieser Plattform. Das ist PaaS.
- Unternehmen D geht einen Schritt weiter und verzichtet auf eine eigene Abacus-Installation. Der Abacus-Mandant ist bei einem Abacus-Vertriebspartner gehostet, und der Zugriff auf die Buchhaltung erfolgt über den Web-Browser. Ein direkter Serverzugang oder der Zugriff auf das Backend ist nicht vorgesehen, weil auf derselben physischen Infrastruktur auch weitere SaaS-Kunden betrieben werden. Das ist SaaS in einer Public-Cloud-Umgebung.

¹² In der Praxis ist ein Backup ausserhalb des eigenen Serverraums erforderlich, da bei Ereignissen wie einem Brand sonst sämtliche Daten verloren gehen könnten. Wird ein automatisierter Backup-Prozess über das Internet bei einem Drittanbieter eingerichtet, wird damit die reine On-Premise-Architektur bereits verlassen.

¹¹ Eigene Darstellung; Quelle: Zacker, Craig (2020), S. 13 ff.

1 Grundlagen: Daten- und Datenbanken

Warum sich die Cloud durchgesetzt hat, zeigen drei zentrale Vorteile:

- *Kosten:* Cloud-Modelle sind oft deutlich günstiger. Viele Unternehmen wollen oder können keine spezialisierten IT-Mitarbeitenden beschäftigen, die im Serverraum Hardware warten, manuell Updates einspielen und Sicherheitsdienste betreiben. In der Cloud übernimmt der Anbieter diese Aufgaben, skaliert Ressourcen und verteilt Kosten.
- *Performance/Skalierung:* Braucht ein Prozess einmal im Monat die zehnfache Rechenleistung, müsste On-Premise auf diese Leistungsspitze hochdimensioniert werden, was hohe Kosten verursacht und 95 Prozent der Zeit überdimensioniert ist. In der Cloud stehen grosse, gemeinsam genutzte Ressourcen bereit, die Lastspitzen abfedern.
- *Sicherheit:* Professionelle Cloud-Anbieter verfügen über anerkannte Zertifizierungen und spezialisierte Security-Teams, die globale Angriffe abwehren. Für ein Kleinunternehmen ist dieses Niveau schwer zu erreichen.

Nicht zu übersehen ist der Nachteil, dass mit der Auslagerung ein Teil der Kontrolle abgegeben wird. Das gilt für jedes Outsourcing. Wichtig ist deshalb ein internes Monitoring, verbunden mit Grundwissen und Vokabular – genau das will diese Broschüre vermitteln.

Wer der Cloud grundsätzlich misstraut, kann den Blick auf ein vertrautes Beispiel richten. Besitzen Sie privat E-Banking? Dann beantworten Sie folgende beiden Fragen:

-
- Gehört die Hardware, auf der Ihre Bankdaten liegen, nicht Ihnen, sondern der Bank?
-
- Bearbeiten Sie Ihre Daten über eine Verbindung via Internet?

Wenn Sie beide Fragen mit Ja beantworten, ist E-Banking ein Cloud-Dienst. Und die Bank stellt nicht für jeden Kunden einen eigenen Server mit separatem Betriebssystem bereit. Mehrere Kundendaten liegen mandantengetrennt auf gemeinsamer Infrastruktur. Das entspricht der Public-Cloud-Definition. Dass solche Daten besonders schützenswert sind, liegt auf der Hand. Entsprechend setzen Banken und SaaS-ERP-Anbieter vergleichbare Schutzmechanismen ein. Aus Endanwendersicht ist die Multifaktor-Authentifizierung der sichtbarste Mechanismus.

Data Analytics und Business Intelligence

Fabian Meisser



2.1 Definitionen und Reifegrade

In der Literatur wird die Datenanalyse oft nach Reifegraden unterteilt. Die vier gängigen Stufen sind **deskriptiv**, **diagnostisch**, **prädiktiv** und **präskriptiv**:

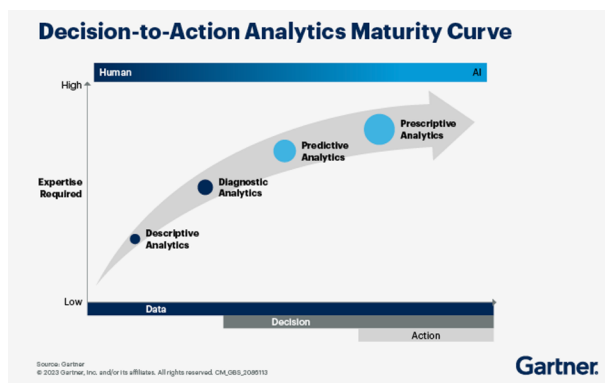


Abbildung 14: Reifegrade von Datenanalyse¹³

Die **deskriptive** Analyse beschreibt die Vergangenheit und macht Abweichungen sichtbar, etwa in einer Erfolgsrechnungsanalyse mit Vorjahresvergleich im Geschäftsbericht. Die **diagnostische** Analyse geht einen Schritt weiter, erklärt Abweichungen datenbasiert und leitet die Treiber systematisch her, statt nur übergeordnete Gründe anzuführen. Die **prädiktive** Analyse (Predictive Analytics) richtet den Blick nach vorne und schätzt, wie sich Kennzahlen entwickeln könnten, gestützt auf historische Muster und ergänzt um Annahmen zu Einflussfaktoren wie Nachfrage, Preisen oder Kapazität. Die **präskriptive** Analyse schliesst daran an und empfiehlt konkrete Massnahmen, die sich aus der Vorausschau ergeben; ein anschauliches Beispiel stammt aus dem Maschinenunterhalt: Wenn die Analyse einen Ausfall wegen veralteter Teile prognostiziert, lautet die Handlungsempfehlung klar und terminlich präzise, Ersatzteil X bis Ende Q1 2028 ersetzen. Diese Abfolge hilft Accountants und Controllern, Berichte nicht nur zu beschreiben, sondern Ursachen zu verstehen, Zukunft zu antizipieren und daraus zielgerichtete Entscheidungen abzuleiten.

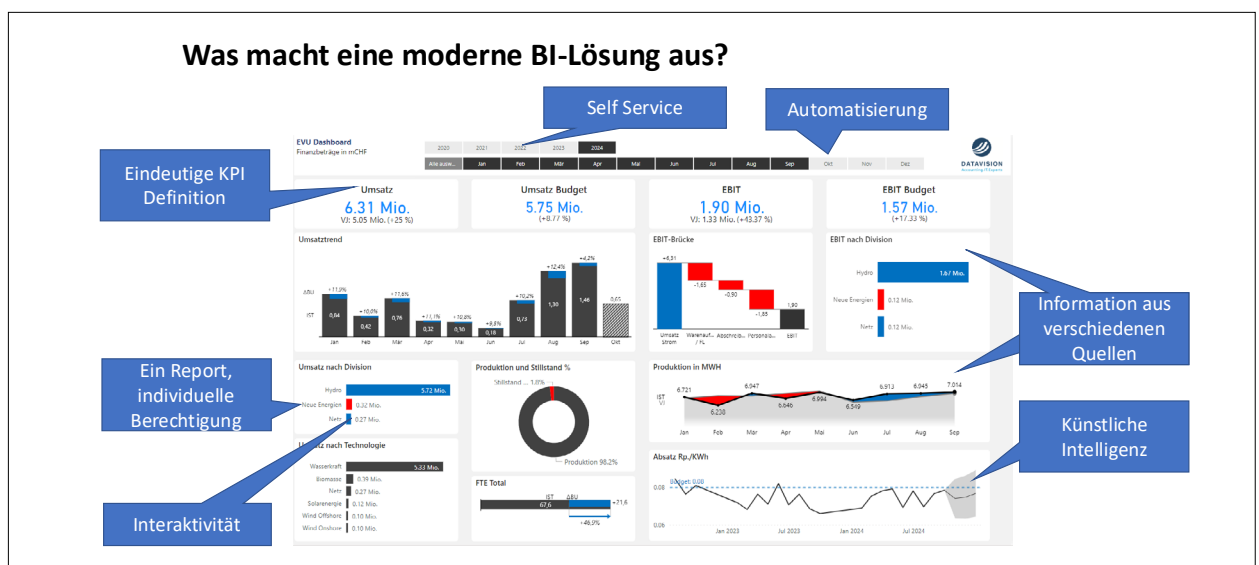


Abbildung 15: Charakterisierung moderner BI-Systeme¹⁴

2 Data Analytics und Business Intelligence

Business Intelligence ist ein umfassendes digitales System zur Automatisierung der Analyse von Unternehmensdaten. Das lateinische Wort «intelligere» bedeutet «verstehen». Es geht also darum, das Unternehmen mit Hilfe von Daten zu verstehen und weniger um «Intelligenz» im Sinn eines Intelligenztests.

1. Zentrale Kennzahldefinition

Durch die Anbindung der Vorsysteme und die dadurch ermöglichte Vollautomatisierung werden Kennzahlen zentral definiert. Im Gegensatz zur manuellen Reporterstellung, bei der in Excel oft noch letzte Anpassungen erfolgen und verschiedene Personen zu unterschiedlichen Resultaten kommen, sorgt eine zentrale Formel für klare, unternehmensweit gültige Definitionen.

2. Self Service

Moderne BI-Tools ermöglichen es Endnutzern, mit den passenden Berechtigungen selbst auf die benötigten Informationen zuzugreifen. Nicht jede Auswertung muss zuerst durch das Controlling erstellt werden, sondern steht direkt im System bereit.

3. Automatisierung

Die Quellsysteme sind dauerhaft an das BI-System angebunden, welches regelmässig automatisch aktualisiert wird. Typischerweise geschieht dies mindestens einmal pro Nacht, sodass Berichte und Dashboards stets auf dem neuesten Datenstand basieren.

4. Information aus verschiedenen Quellen

Viele Kennzahlen entstehen aus der Kombination mehrerer Systeme. «Umsatz pro Mitarbeitende» lässt sich beispielsweise aus dem Kassensystem für den Umsatz und dem HR System für die Anzahl Mitarbeitende ableiten. BI ist damit mehr als die Abbildung eines einzelnen Systems wie Abacus, SAP oder Business Central, sondern eine integrierte Sicht über mehrere Quellen hinweg.

5. Interaktivität

Im Hintergrund arbeitet ein Datenmodell, das flexible Kombinationen von Dimensionen ermög-

licht. Nutzerinnen und Nutzer können interaktiv filtern und analysieren, etwa ein Divisionsergebnis anklicken und das gesamte Dashboard dynamisch auf diese Division einschränken.

6. Individuelle Berechtigungen

Über das persönliche Login sieht jede Person nur die Daten des eigenen Zuständigkeitsbereichs. So können sich beispielsweise 50 Kostenstellenleitende einen Bericht teilen, und beim Login ist die Sicht jeweils auf die eigene Kostenstelle vorgefiltert.

Im Zusammenhang mit der Bereitstellung von Unternehmensdaten gibt es in der Praxis verschiedene Begrifflichkeiten. Früher war «MIS» (Management Information System) geläufig. Der Fokus auf das «Management» als primäre Zielgruppe greift heute jedoch zu kurz. In einem modernen Führungsverständnis überwacht nicht nur das (Top-)Management Kennzahlen. Die Treiber werden letztlich von vielen Personen im Unternehmen gesteuert, und der Datenzugriff richtet sich auf die Ebene, auf der eine Person tatsächlich Einfluss nehmen kann. Verantwortung für Zielgrößen kann nur übernehmen, wer die notwendige Sichtbarkeit der Daten und ihrer Einflussgrößen hat.

Ein weiterer methodischer Begriff ist die «Balanced Scorecard», ein Modellansatz von Robert S. Kaplan und David P. Norton aus den frühen 1990er-Jahren. Er zielt darauf ab, mehrere Perspektiven abzubilden – beispielsweise die finanzielle Sicht, Umweltaspekte sowie HR und Kultur. Diese Perspektiven lassen sich mit Business Intelligence konsistent modellieren und auswerten; teils sind sie als eigenständige Modelle ausgeprägt, werden aber durch BI in einer gemeinsamen Daten- und Reporting-Logik zusammengeführt.

Business Intelligence entstand durch starke Rechenleistung, die breite Verfügbarkeit von Systemen und Daten sowie den schnellen Aufstieg der Cloud. BI bildet sowohl operative als auch strategische Unternehmensdaten ab – von standardisierten Kennzahlen bis zu interaktiven Steuerungsgrößen für die jeweils verantwortlichen Personen.

2 Data Analytics und Business Intelligence

2.2 Bedeutung von Business Intelligence

Seit dem Aufkommen von Clouddiensten in den 2010er-Jahren haben sich mehrere Trends herausgebildet, die Unternehmen herausfordern und auf die Business Intelligence konkrete Antworten bietet.

• Heterogene Systemlandschaft

Unternehmen – auch KMU – betreiben heute Dutzende bis über hundert Softwaresysteme parallel. Jede dieser Anwendungen erzeugt Daten, die für das Unternehmen relevant sein können. Die integrierten Auswertungen der einzelnen Tools sind oft funktional begrenzt und jeweils nur auf das eigene System fokussiert. Weil die Business Logik systemübergreifend wirkt, enden viele Prozesse in Excel Exporten aus verschiedenen Systemen und in einer fehleranfälligen, manuellen Überleitung. BI setzt hier an, integriert Quellen, harmonisiert Definitionen und automatisiert die Auswertung.

• Granularere Daten

Sinkende Speicher und Verarbeitungskosten machen sehr feine Detailstufen verfügbar. Statt lediglich Monatssalden pro Kostenstelle und Konto stehen heute Daten auf Belegebene bereit, zum Beispiel auf der Stufe Kreditorenrechnung inklusive Lieferant, Freigabedatum und freigebender Person im Kreditoren-Workflow. Wichtig ist zudem, dass es keine Systembrüche gibt, wenn direkt ein Drilldown möglich ist: Liegen die Kosten über Budget, verlässt man nicht das BI und loggt sich separat im ERP ein, sondern sieht unmittelbar im BI, welche Rechnungen oder Buchungen zum Überschreiten geführt haben – bis hin zu den relevanten Belegdaten. Die Granularität reicht in vielen Anwendungen bis auf die Zeitebene mit Uhrzeitstempel; BI nutzt diese Tiefe für präzisere Analysen und belastbare Treiberlogiken.

• VUCA¹⁵-Umfeld

Die Unternehmensumwelt ist volatil, unsicher, komplex und mehrdeutig. Diese Dynamik lässt sich an praktischen Beispielen im eigenen Unternehmen nachzeichnen. Gefragt ist Flexibilität – moderne BI-Tools sind darauf ausgelegt und unterstützen

schnelle Anpassungen von Kennzahlen, Sichten und Berichten bei sich verändernden Anforderungen.

• Bedarf an jederzeit aktuellen Informationen

«Continuous Accounting» steht für die Abkehr von Zwischenabschlüssen hin zu laufend aktualisierten Management-Informationen. Die zentrale Frage lautet: Wo stehen die wichtigsten Treiber heute – oder mindestens bis gestern Abend? Dafür braucht es keinen vollständigen buchhalterischen Abschluss, sondern tagesaktuelle Kerngrößen. In vielen Branchen ist der Deckungsbeitrag (z. B. DB1, DB2) entscheidend: Umsätze minus direkt verbundene Kosten sind der Haupttreiber; Fixkosten sind oft über Verträge im Voraus bekannt und müssen nicht täglich verifiziert werden. Viele Unternehmen erstellen deshalb nur die gesetzlich oder regulatorisch geforderten Abschlüsse und stützen den täglichen Überblick auf die leistungsfähigen Auswertungsmöglichkeiten moderner BI-Tools.

Kurzfazit: BI beantwortet diese Trends mit zentralen Definitionen, integrierten Datenmodellen, automatisierten Ladeprozessen, Drilldowns ohne Systembruch und interaktiven Auswertungen – und schafft damit eine verlässliche, flexible Grundlage für die finanzielle Steuerung im Alltag.

2.3 Das BI-Framework

Welche Schritte sind nun konkret notwendig, um von den Unternehmensdaten auf das Endresultat (ein Dashboard, Standardreport, automatisierte Monatsreports) zu gelangen?

Es wird hier mit folgendem BI-Framework gearbeitet:

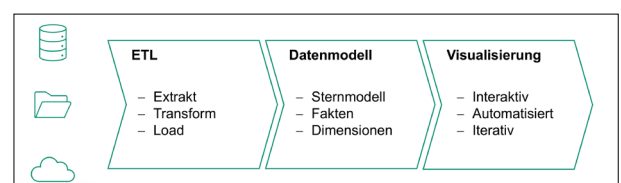


Abbildung 16: Das BI-Framework¹⁶

¹⁵ Volatility, Uncertainty, Complexity, Ambiguity (Mehrdeutigkeit).

Quelle: <https://wirtschaftslexikon.gabler.de/definition/vuca-119684>

¹⁶ DataVison AG (2026).

2 Data Analytics und Business Intelligence

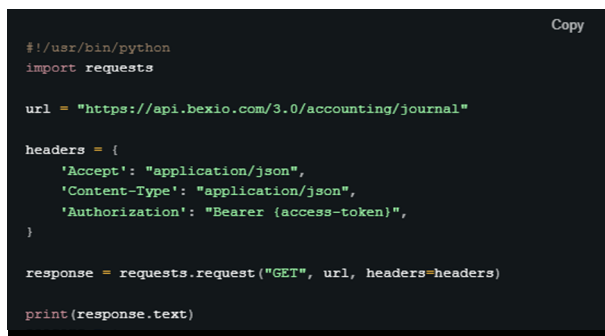
ETL

Im ETL-Prozess (Extract, Transform, Load) werden die als relevant identifizierten Unternehmensdaten, die in unterschiedlicher Form vorliegen können, extrahiert, transformiert und in das BI-Modell geladen.

Beim Extrahieren («E» von ETL) gibt es diverse technische Zugangsformen, deren wichtigste hier erläutert werden:

- Produktspezifische Konnektoren: Vorgefertigte Anschlüsse für ein bestimmtes Produkt, z.B. ein SAP-HANA-Konnektor.
- Generische Konnektoren: Schnittstellen wie ODBC, über die sich verschiedene Produkte wie Abacus oder Sage/Infoniqa anbinden lassen; meist ist die Installation eines passenden Treibers notwendig.
- API-Zugriff: API steht für Application Programming Interface. Datenpakete werden über das Internet bezogen, der verbreitetste Standard ist die REST-API (Representational State Transfer). Eine typische Rückgabe kann als CSV oder JSON vorliegen und muss in ein Zeilen-/Spaltenformat überführt werden.¹⁷

Bei REST werden Ressourcen über Endpunkte angesprochen. Wenn z.B. alle Buchungsjournal-einträge bezogen werden sollen, führt ein Endpunkt wie <https://api.bexio.com/2.0/accounting/journal> zum Ziel:



```
#!/usr/bin/python
import requests

url = "https://api.bexio.com/3.0/accounting/journal"

headers = {
    'Accept': "application/json",
    'Content-Type': "application/json",
    'Authorization': "Bearer {access-token}",
}

response = requests.request("GET", url, headers=headers)

print(response.text)
```

Abbildung 17: API-Abfrage bei Bexio¹⁸

¹⁷Ein öffentliches CSV-Beispiel lässt sich hier finden: <https://raw.githubusercontent.com/owid/covid-19-data/master/public/data/owid-covid-data.csv>

Organisatorisch gilt: Neben dem technischen Zugang braucht es einen Datenbank- oder API-Benutzer mit den notwendigen Berechtigungen. Sobald ein BI- oder Analytics-Tool über einen Konnektor zugreift, folgt die Authentifizierung, je nach System mit Benutzername und Passwort, API-Key, OAuth 2.0 oder einem Service-Konto.

Der wichtigste Buchstabe in ETL ist wohl das T (für Transformation). Wenn man davon ausgeht, dass diese Daten korrekt sind, wieso sollen sie dann noch transformiert werden? Ist hier nicht die Gefahr, die Daten zu verfälschen? Nun, die Gründe sind vielfältig, ein paar Beispiele für übliche Transformationen:

- **Auswahl von Datenfeldern:** Eine Tabelle in Rohdatenform hat oft unzählige Spalten, die technisch in der Datenbank mitgeführt werden, aber für Auswertungszwecke nicht gebraucht werden. Z.B. die Buchungstransaktionstabelle von Abacus besteht aus über 180 Spalten, davon sind in einem typischen BI Setup rund 20 bis 30 relevant.
- **Filterung von ungültigen Werten:** Oft sind auch stornierte Werte oder Testwerte in den Tabellen enthalten, die zu falschen Saldi führen würden. Oder ein Unternehmen führt im selben System noch eine Mandatsbuchhaltung, deren Daten nicht in den Report der eigenverantwortlichen Einheit einfließen sollen – all dies sind Beispiele für Filtervorgänge.
- **Zusammenführen von Daten** («Join», «Merge»): Das vereinfachte Pendant dazu in Excel wäre ein SVERWEIS oder XVERWEIS (VLOOKUP, XLOOKUP). Über einen gemeinsamen Schlüssel werden Daten aus zwei Tabellen in eine zusammengeführt. Der Grund dafür ist, dass transaktionale Systeme wie ein ERP darauf ausgerichtet sind, die Daten möglichst optimal fürs interne Funktionieren des Systems zu speichern. So sind Informationen zu einem Thema (z.B. Kundenstammdaten) oft auf viele Tabellen verteilt. Für die Auswertung gehören diese Informationen jedoch zusammen.

¹⁸<https://docs.bexio.com/#tag/Reports/operation/ListJournalEntries>

2 Data Analytics und Business Intelligence

• Bereinigen von fehlenden oder unvollständigen

Datensätzen: Fehlerhafte Werte kommen in der Theorie kaum vor, in der Praxis jedoch zuhauf. Ein Beispiel zum Selberausprobieren: Geben Sie in Excel 0/0 ein und laden Sie die Tabelle in ein BI-System. Wenn man weiss, dass solche Fehler fachlich vernachlässigbar sind und mit «n/a» ersetzt werden können, setzt man eine Regel, damit nicht deswegen der ganze Prozess abbricht.

Beim Laden («L» von ETL) werden die transformierten Daten ins Modell übernommen und gemäss Zeitplan automatisch aktualisiert, typischerweise mindestens einmal pro Nacht. Für höhere Aktualität kann die Frequenz erhöht werden.

Datenmodell

Wenn mehrere Tabellen in ein BI-System geladen werden, können sie nur gemeinsam ausgewertet werden, wenn ihre Beziehungen klar sind. Als Best Practice in BI-Projekten gilt das Sternenmodell oder englisch: Star-Modell. Dabei unterscheidet man zwischen Faktentabellen (transaktionale Daten) und Dimensionen (Stammdaten). Damit diese verbunden werden können, braucht es gemeinsame Schlüssel («Keys»). Ein klassisches Beispiel ist die Kundennummer. Für ein stabiles System ist es wichtig, stets einen Schlüssel und nie einen Personennamen als Schlüssel zu verwenden, denn Kundennamen können sich ändern – bei Privatpersonen z.B. durch Heirat, bei Unternehmen durch Umfirmierungen.

Beziehungen haben weitere Unterausprägungen:

- **Kardinalität:** Schlüssel können in einer Tabelle eindeutig («unique», also nur einmal vorkommend) oder mehrfach vorkommen. Für «mehrfach» wird oft «N» als Abkürzung gebraucht. Wenn in der Kundenstammtabelle der Kunde mit Nr. 721 genau einmal vorkommt und in der Transaktionstabelle «Verkäufe» mehrfach, ist die Beziehung 1:n. Sind beide Tabellen mehrdeutig, spricht man von M:n.
- **Kreuzfilterrichtung:** Beziehungen können einseitig sein (Standardfall: Die Dimension filtert die Faktentabelle) oder beidseitig (beide Tabellen filtern

sich gegenseitig). Gewisse BI-Tools lassen nur 1:n-Beziehungen und einfache Filterrichtungen zu. Diese Kombination ist Best Practice und Abweichungen sollten nur in begründeten Fällen oder mit einer sinnvollen Alternative in Betracht gezogen werden.

Folgende Anmerkungen zur praktischen Ausgestaltung eines Datenmodells:

- Ein Datenmodell hat typischerweise mehrere Faktentabellen, nicht nur eine wie in den obenstehenden, vereinfachten Grafiken. Neben den Fakten (z.B. effektive Umsätze) ist im Finanzkontext das Budget eine typische weitere Faktentabelle. Sie enthält deutlich weniger Details und benötigt andere Spaltenstrukturen als eine Ist-Tabelle. Beispiel: Im Budget gibt es im Gegensatz zum Ist keine Buchungs-, Kreditoren- oder Lieferantenummern, sondern meist einen Saldobetrag pro Monat, Konto und Kostenstelle. Eine weitere Faktentabelle kann auch eine eher qualitative Ergänzung sein, z.B. die Anzahl und Art von Kundenreklamationen.
- Die Elemente eines Datenmodells unterscheiden sich von Unternehmen zu Unternehmen. Ein Spital könnte Dimensionen haben wie Tarifmodell, Patientinnen und Patienten, Stationen und Behandlungen – Begriffe, die bei einem Generalunternehmer keinen Sinn ergeben würden. Gleichzeitig bildet eine Organisation meist nicht die gesamte Zahlenwelt in einem einzigen Modell ab, weil die Komplexität zu gross würde. In der Fallstudie in diesem Kapitel wird anhand einer Gastronomiegruppe aufgezeigt, wie sich dies konkret auswirken kann.

Definition von Kennzahlen

Sobald ein Datenmodell vorliegt, ist ein typischer nächster Schritt die Definition von Kennzahlen, oft auch «Measures» genannt. Bei den BI-Tools in Excel und bei Power BI ist die Sprache dafür DAX (Data Analysis Expressions). Es handelt sich um eine Formelsprache, die auf dem Datenmodell und dessen Eigenheiten aufsetzt. Ein wichtiger Aspekt dabei ist, dass eine Kennzahl nicht wie in Excel in einer einzel-

2 Data Analytics und Business Intelligence

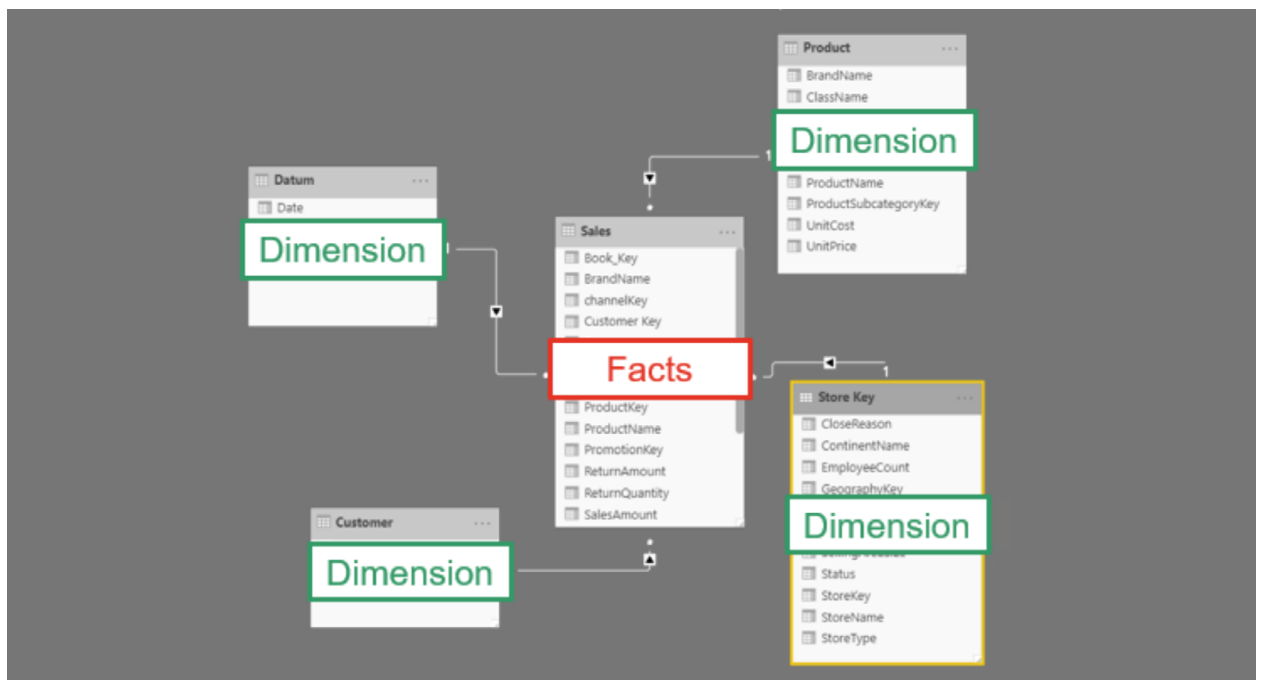


Abbildung 18: Star Schema eines BI-Modells (eigene Darstellung)

nen Zelle «gefangen» ist, sondern modellübergreifend funktioniert und auf neuen Berichtsseiten wiederverwendet werden kann. Das macht das Ganze stabiler und mächtiger, aber auch weniger flexibel als eine frei platzierte Excel-Formel.

Die Syntax einer beispielhaften Kennzahl ist dabei wie folgt:

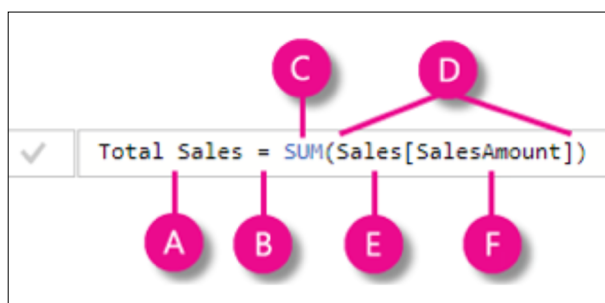


Abbildung 19: Syntax von DAX (eigene Darstellung)

- A. Name des Measures (Kennzahl)
- B. Gleichheitszeichen (=), hier beginnt die Formel
- C. DAX – Funktion (hier Summe)
- D. Klammer für Summenfunktion
- E. Relevante Tabelle
- F. Relevante Spalte («Attribut»)

Es gibt über 250 Funktionen, viele davon sind gerade im Finanzkontext sehr wertvoll. Beispiele:

- Basisfunktionen wie SUM, DIVIDE etc.
- Zeitintelligenz wie YTD, MTD, TODAY etc.
- Finanzmathematische Funktionen wie FV, XNPV etc.

Eine Sonderrolle nimmt die sehr wichtige Funktion CALCULATE ein. Diese Funktion steuert respektive übersteuert den aktuellen Filterkontext (also z. B. ein Filter, welches zuoberst in einem Dashboard das Jahr einschränkt). Wenn also z. B. der Filter in einem Dashboard auf 2026 gesetzt ist und man trotzdem dynamisch den Vorjahreswert (2025) anzeigen möchte, ist CALCULATE die Antwort. Ebenso bei Feinjustierungen: Der adjustierte EBIT soll der bereits vorhandenen Kennzahl «EBIT» entsprechen, aber um beispielsweise drei Konten gemäss Kontenplan Dimension bereinigt sein.

2 Data Analytics und Business Intelligence

Visualisierung

Zu BI hört man mitunter die Aussage «BI ist nur ein Visualisierungstool». Dies gilt jedoch nicht für moderne BI-Tools, welche alle Schritte des hier verwendeten Frameworks End to End abbilden können.

Wie man eine Visualisierung aus einem Datenmodell erstellt, ist meist sehr toolspezifisch und intuitiv. Eine Programmierung ist in den führenden Tools nicht oder nur ausnahmsweise notwendig. Spannender sind die Best Practices in diesem Zusammenhang.

Ähnlich der Standardsetzung im Accounting gibt es zwei unterschiedliche Ansätze:

Der prinzipienbasierte Ansatz gibt gewisse Grundsätze vor, ohne die kleinsten Details zu regeln. Das «ACTION»-Akronym ist ein solches Beispiel, welches die gängigsten Best-Practice-Regeln im Zusammenhang mit Dashboards zusammenfasst. Es ist quasi der kleinste gemeinsame Nenner, auf den sich die Praktiker in 90 % der Fälle einigen könnten:

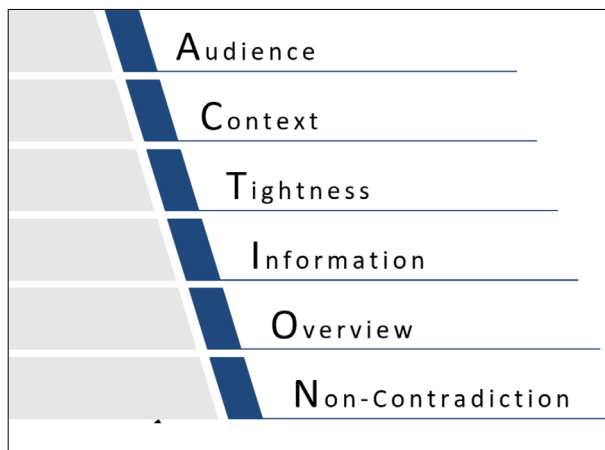


Abbildung 20: Das Action-Akronym für Dashboards¹⁹

- **Audience:** Eine Visualisierung muss zielgruppen-gerecht aufgebaut werden. VR-Dashboards haben andere Anforderungen als ein Bericht für Kostenstellenleiterinnen.
- **Context:** Nackte Kennzahlen (EBIT = CHF 5 Mio.) bringen wenig Mehrwert, wenn nicht eine Kontextgrösse mitgegeben wird. Typische Kontextgrössen sind Budget, Vorjahr, Zielwerte oder Branchen-Benchmarks.
- **Tightness:** Daten sollten zuerst verdichtet gezeigt werden. Statt die Auflistung von 100 Landesgesellschaften, zunächst die vier Verkaufshauptregionen.
- **Information:** Kennzahlen, Achsen und Titel sollen sinnvoll beschriftet werden; technische, kryptische Begriffe vermeiden, wenn die Zielgruppe diese voraussichtlich nicht versteht.
- **Overview:** Man kann bei einem Dashboard an die Analogie zum Armaturenbrett im Auto denken: Es sollte aufgeräumt und übersichtlich sein.
- **Non-contradiction:** Die einzelnen Aussagen sollten sich nicht widersprechen – einerseits inhaltlich (oben links «über Budget», unten rechts «unter Budget» wäre widersprüchlich), andererseits bei Farbsignalen. Farben sollen eine Bedeutung haben und bewusst gewählt werden. Zwei unterschiedliche Dunkelblau-Töne (einmal für ein Land, einmal für ein Produkt) verwirren die Berichtsempfänger.

Der regelbasierte Ansatz hat zum Ziel, sehr detailliert Regeln vorzugeben, wie mit Visualisierung zu verfahren ist. Ein Beispiel ist der International Business Communication Standards (IBCS).²⁰ Hier gibt es ein Regelwerk mit diversen Normen, wie Istwerte ausgefüllt, die Planwerte leer und die Forecasts schräg schraffiert. Dies hat zum Ziel, dass sich die Berichtsempfänger mit der Zeit nicht mehr an Legenden orientieren müssen, sondern dass weltweit für alle klar ist, welche Darstellungsform was bedeutet.

¹⁹ DataVision AG (2026)

²⁰ www.ibcs.com/de Über diesen Link kann eine kostenlose Broschüre mit empfohlenen Notationen heruntergeladen werden.)

2 Data Analytics und Business Intelligence

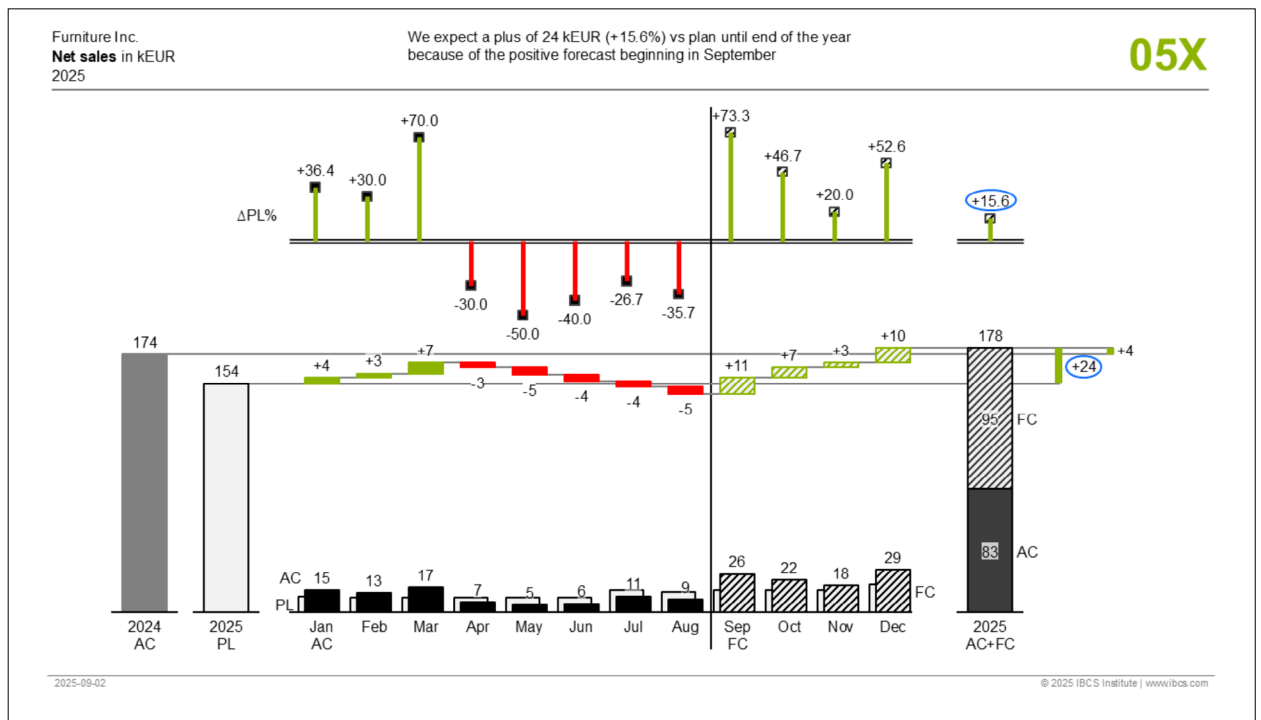


Abbildung 21: Reporting nach ISCB-Standards²¹

2.4 Formen des finalen Reportings

Digitale Berichte bieten Möglichkeiten, die sie um ein Vielfaches mächtiger machen als ausgedruckte Berichte. Man stelle sich folgendes, realistisches Szenario vor: Zahlen werden in einer Sitzung an die Wand projiziert, und der CEO hinterfragt eine Kennzahl («Das kann doch nicht stimmen»). In der klassischen Situation ohne BI suchen nun verschiedene Personen nach einer Erklärung, können diese aber nicht datengestützt direkt in der Sitzung untermauern. Ein starkes Controlling im Zeitalter von Big Data klickt auf die Zahl und erläutert sofort die Hintergründe. Hierbei gibt es beispielsweise folgende Untervarianten:

- **Drilldown** (innerhalb derselben Visualisierung schrittweise in einer Hierarchie von der aggregierten Ebene in detailliertere Daten wechseln; z.B. Jahr → Quartal → Monat).
- **Drillthrough** (auf eine andere Berichtssseite springen, die automatisch anhand der Auswahl als

Kontextfilter vorgefiltert ist und Detailansichten zur gewählten Kennzahl zeigt).

Quickinfo (kontextabhängige Anzeige beim Überfahren mit der Maus)

Varianzerklärung, z. B. mit Wasserfalldiagramm-Varianten (Zerlegung der Veränderung zwischen zwei Werten – z. B. Plan vs. Ist, Vorjahr vs. laufendes Jahr – in beitragende Faktoren).

Ist BI geeignet für ausdrückbare Berichte? Ja, wenn die passende Form gewählt wird. Interaktive BI-Berichte sind zum Erkunden gedacht: Filter und Drilldowns verändern den Umfang, Tabellen wachsen, und der Bildschirm zeigt jeweils nur den aktuellen Ausschnitt. Für Anhänge und fixe Auszüge eignen sich paginierte Berichte mit festen Seiten, wiederholten Kopfzeilen und pixelgenauer Darstellung. Der PDF-Export bleibt stabil. Ein typischer Einsatz ist beispielsweise eine tabellarische Liste aller Immobilien oder Projekte mit ihren Kenngrößen.

Beim Zugriff auf Inhalte sind zwei Prinzipien zentral. Row Level Security (RLS) filtert Daten auf Zeilen-

²¹ISBC.com (2025), <https://www.ibcs.com/de/resource/horizontal-waterfall-chart/>

2 Data Analytics und Business Intelligence

ebene je Rolle, damit jede Person nur das sieht, was sie betrifft, zum Beispiel die eigene Kostenstelle oder Region. Die Rollen werden im Datenmodell definiert und können dynamisch über die Anmelde-E-Mail gesteuert werden. So bleibt die Sicht konsistent und testbar. Role Based Access Control (RBAC) legt fest, welche Aktionen erlaubt sind. Arbeitsbereichsrollen und Datensatzberechtigungen trennen klar zwischen Lesen für die reine Konsumation und erweiterten Rechten wie Erstellen, Teilen oder Ändern. Das schafft nachvollziehbare Verantwortlichkeiten und vereinfacht Audits.

In der Praxis gibt es oft die Situation, dass Verantwortliche aufgefordert werden, einmal pro Monat einen Bericht anzusehen. Das verpufft leicht im Tagesgeschäft. Wirkungsvoller ist eine Push-Verteilung der Berichte: Relevante Informationen kommen automatisch zur richtigen Zeit. Eine wöchentliche E-Mail mit einer kompakten Übersicht zu offenen Debitoren markiert Überfälligkeiten klar und verlinkt direkt in den Bericht. Bei der Budgetkontrolle macht eine kurze Benachrichtigung Abweichungen sofort sichtbar und führt mit einem Klick zur betroffenen Kostenstelle. So entsteht kein zusätzlicher Erinnerungsaufwand, und die Aufmerksamkeit der Verantwortlichen wird dorthin gelenkt, wo sie gebraucht wird.

2.5 Toolübersicht

Welche Tools eignen sich für ein Business-Intelligence-Vorhaben? Die Toollandschaft verändert sich schnell, ebenso wie die Unternehmensumwelt. Gleichzeitig bedeutet die Einführung eines neuen Tools spürbaren Aufwand, sodass ein jährlicher Wechsel nicht opportun ist – vergleichbar mit der Einführung eines ERP. Hilfreich sind neutrale Vergleichsdienste, die ein Themengebiet über Jahre hinweg beobachten und regelmässig gegenüberstellen. Ein Beispiel ist Gartner: Im Bereich «Analytics and Business Intelligence Platforms» lässt sich über mehr als zehn Jahre nachverfolgen, welche Lösungen als führend eingestuft wurden.

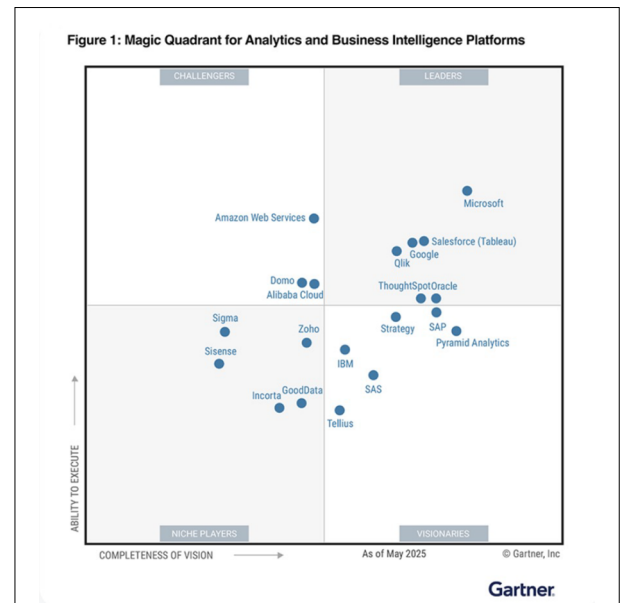


Abbildung 22: Gartner-Vergleich von BI-Systemen²²

In der Mehrjahresanalyse 2015–2025 erscheinen als führende Anbieter immer wieder dieselben drei Unternehmen: Microsoft (Power BI), Salesforce (Tableau) und Qlik (QlikView/Qlik Sense). Diese Lösungen sind auf systemübergreifende Auswertungen ausgelegt; andere Tools fokussieren stärker auf einzelne Quellsysteme, etwa ERP-integrierte BI-Anwendungen wie SAP Analytics Cloud (SAC) oder der Data Analyzer in Abacus. Microsoft verfolgt bewusst einen offenen Ansatz und stellt in Power BI über 200 Datenkonnektoren bereit – auch zu Konkurrenzsystemen wie SAP, IBM oder Oracle-Datenbanken. In Excel sind mit Power Query (fest integriert) und Power Pivot (als Add-in zu aktivieren) die Grundlagen für Power BI bereits gelegt. Beide Funktionen lassen sich parallel nutzen und bieten einen besonders einfachen Einstieg in die BI-Welt; in vielen Unternehmen erfordert Power Query dafür weder zusätzliche Freigaben noch ein separates Budget.

²²Gartner, 2025

2 Data Analytics und Business Intelligence

Case Study: Datenbanken, ERP-Systeme und Business Intelligence in einem Immobilienunternehmen

Das Immobilienunternehmen Immo AG mit Sitz in Zürich hält und bewirtschaftet ein Portfolio von 50 Liegenschaften. Neben der Kerntätigkeit baut die Firma eine eigene Sparte Generalunternehmer (GU) auf. Für die Bewirtschaftung ist eine operative Software im Einsatz, die Mietabrechnungen, kleinere Unterhaltsarbeiten, Mieterkommunikation und Mieterwechsel zuverlässig abwickelt. Für die strategische Steuerung des Portfolios, inklusive Simulationen von langfristigen Erweiterungsbauten, wurde zusätzlich ein passendes Portfolio Tool eingeführt.

Beide Systeme bieten jedoch keine Lohnbuchhaltung. Dadurch stellt sich die Frage nach der optimalen Gesamtlösung für Finance und Payroll, gerade mit Blick auf die neue GU-Sparte.

Der ERP-Anbieter, über den neu die Buchhaltung laufen soll, bietet zwar ein Bau-Modul an. Die Evaluation zeigt jedoch, dass dieses eher für kleine, weniger komplexe Bauvorhaben geeignet ist. Immo AG plant als Generalunternehmer Grossprojekte mit bis zu 100 Mio. CHF Bauvolumen und Dutzenden Subunternehmen. Im Austausch mit anderen GU wird klar, dass es eine über die letzten Jahre optimierte Spezialsoftware gibt, die genau auf Grossprojekte mit schweizerischen Spezifika wie Regulierungen und Bewilligungen ausgerichtet ist. Immo AG entscheidet sich deshalb für die Einführung dieses GU-Tools in Kombination mit dem ERP für die Finanzprozesse inkl. Payroll.

Da die Gesamtkosten aller Tools höher sind als das frühere, einfache Setup, werden Cloud-Varianten geprüft. Die SaaS-Variante in der Public Cloud erweist sich im Betrieb als rund 50 % günstiger als der Eigenbetrieb.

Abgerundet wird die Softwarestrategie, indem die Daten aus allen Systemen zentral in ein BI-System integriert werden. Dadurch entfällt die Notwendigkeit, die einzelnen Analytics Module der jeweiligen Produkte separat zu lizenzieren, und Auswertungen können unternehmensweit konsistent bereitgestellt werden.



Jon Cajacob

3.1 Überblick

In diesem Kapitel erweitern wir unseren Blick auf alle Datenbestände und Datenflüsse einer Organisation. Dieses Gesamtkonstrukt nennen wir Datenarchitektur. Es lohnt sich, dies zu vertiefen, denn heute ist es die Norm, dass wir nicht nur eine isolierte Softwareanwendung, sondern eine Vielfalt an Applikationen und Systemen brauchen, um die Prozesse des Unternehmens zu unterstützen.

Nehmen wir das Beispiel einer Auftragsfertigung: Ein neuer Kunde ruft an und möchte eine Bestellung machen. Dafür wird er im CRM erfasst und es wird eine Offerte erstellt. Dann wird im ERP ein Kundenauftrag angelegt, welcher ein Produktionsauftrag auslöst. Die Mitarbeiter entnehmen dafür Materialien aus dem Lager und verbuchen ihre Arbeitszeit

im HR-Tool. Nach Auslieferung des Produkts wird eine Rechnung erstellt und in der FIBU verbucht.

Für einen solchen Ablauf brauchen wir bereits mehrere Systeme, in denen laufend neue Daten erfasst oder bestehende verändert werden. Und natürlich müssen diese Applikationen miteinander «sprechen» können – über entsprechende Schnittstellen.

Wieso nun aber ist das Thema Datenarchitektur von Relevanz? Mit einer robusten, integrierten Datenarchitektur bilden wir die bestmögliche Basis nicht nur für operative Systeme, sondern auch für analytische Anwendungen aller Art. Sei es ein BI-Dashboard oder eine KI-Anwendung: Wir wollen diese ordentlich in die Systemlandschaft einbetten, Daten von bestmöglicher Qualität als Input verwenden, die Risiken (z.B. Cybersecurity) jederzeit im Griff haben und dabei effizient bleiben und Doppelspurigkeiten vermeiden.

Die folgende Darstellung zeigt eine typische Datenarchitektur in einer Organisation:

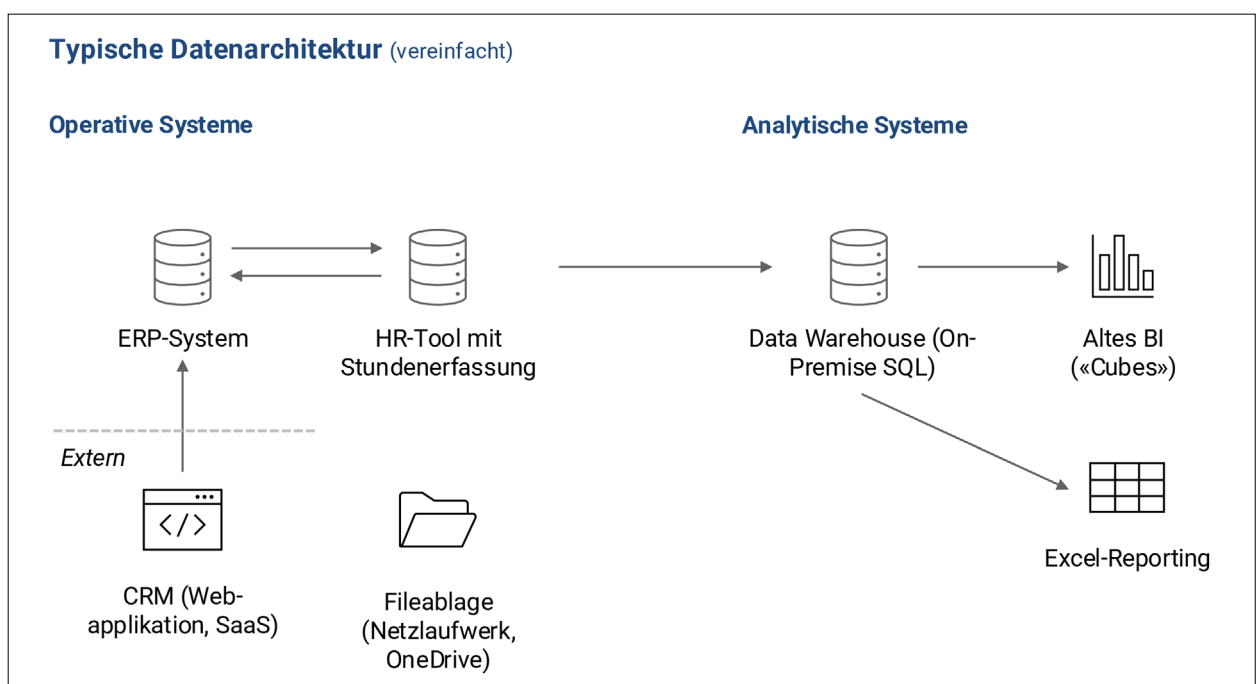


Abbildung 23: Vereinfachte Darstellung einer typischen, über die Zeit gewachsenen Datenarchitektur (eigene Darstellung)

3 Datenarchitektur

Anhand dieser Abbildung können wir einige wesentliche Merkmale der Bestandteile einer Architektur definieren.

Operative vs. analytische Daten

Operative Anwendungen unterstützen die Prozesse des Tagesgeschäfts unmittelbar. Daten werden laufend erfasst und verändert. Die entsprechenden Applikationen sind kritisch für den ordentlichen Geschäftsbetrieb und ein möglichst fehlerfreier Betrieb ist deshalb wichtig.

Analytische Anwendungen, z.B. das BI-Tool, unterstützen die Entscheidungsfindung. Datenbestände werden meist einmal täglich aktualisiert und aufbereitet (z.B. in einem Dashboard). Man unterscheidet zwischen der beschreibenden (was ist passiert) und der voraussagenden (was wird passieren) Analyse. Daten für analytische Tools kommen grösstenteils aus operativen Systemen.

Stammdaten und Bewegungsdaten

Stamm- und Bewegungsdaten (oder Faktendaten) wurden bereits in Kapitel 1.2 beschrieben. Im Kontext der Datenarchitektur ist dies ein wichtiges Unterscheidungsmerkmal, da sich Stammdaten eher langsam und Bewegungsdaten hingegen eher schnell verändern. Oftmals sind Stammdaten auch applikationsübergreifend relevant und sollten entsprechend synchron gehalten werden.

Interne vs. externe Datenbestände

Interne Datenbestände sind jederzeit im eigenen Netzwerk gespeichert (z.B. lokale Datenbank). Früher war es die Norm, dass man fast ausschliesslich interne Datenbestände hatte. Die heutige Popularität von sogenannten SaaS-Tools (Software-as-a-Service) wirkt dem aber entgegen und unternehmensexterne Datenbestände nehmen entsprechend zu.²³ Hierbei sind die Daten bei einem Softwareanbieter, also einer Drittpartei, gespeichert und der Datenzugriff erfolgt über das Internet (bzw. Cloud).

Viele Unternehmen unterschätzen dann auch tendenziell das Risiko von externen Datenbeständen, denn das Internalisieren (z.B. Backups machen) der eigenen Daten wird oftmals technisch nicht vom externen Softwareanbieter unterstützt (oder nur gegen Aufpreis).

Datenaktualität

Wie oft ein Datenbestand aktualisiert wird, ist ein wichtiges Thema, denn die technische Komplexität steigt zusammen mit der Intensität der Datenaktualisierung. Eine tägliche Aktualisierung (Batch-Daten) ist weitaus simpler und weniger fehleranfällig als «Real-Time» Datenströme. Hierbei gilt es einen guten Mittelweg zu finden pro System und Prozess.

Zentrale vs. dezentrale Organisation

Meist wird die Datenarchitektur von zentraler Stelle aus in der Organisation definiert und implementiert. Dies ist insbesondere für operative Datenbestände fast ausschliesslich der Fall. Im analytischen Bereich jedoch gibt es zunehmend einen Trend in Richtung dezentrale Architektur²⁴, das heisst analytische Anwendungen wie z.B. das BI-Tool werden direkt im Controlling-Team statt durch eine zentrale IT-Stelle erstellt und gepflegt. Ein wesentlicher Treiber davon sind der Self-Service sowie der zunehmende Low-Code Funktionsumfang von modernen BI-Tools.

Konsequenzen einer vernachlässigten Datenarchitektur

Um die Wichtigkeit einer robusten Datenarchitektur zu unterstreichen, soll kurz beispielhaft dargestellt werden, was passiert, wenn wir die Architektur vernachlässigen. Dazu die folgende Abbildung:

²³ Vgl. Reis/Housley (2022), S. 8 ff.

²⁴ Vgl. Dehghani (2022), S. 19 ff.

3 Datenarchitektur

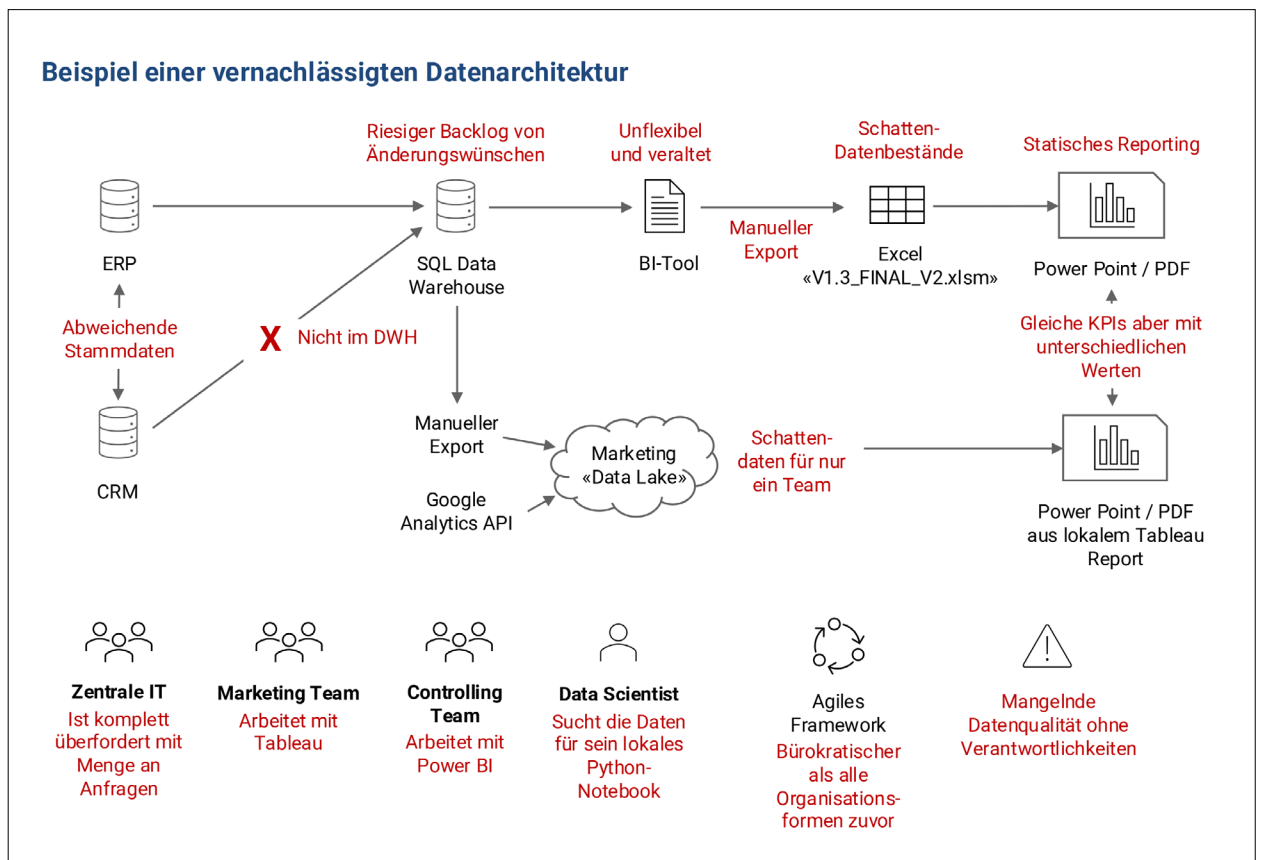


Abbildung 24: Beispiel einer vernachlässigten Datenarchitektur und deren Konsequenzen (eigene Darstellung)

Wenn Sie das anschauen und sich denken: «Das ist ja übertrieben, so wird es in der Realität kaum sein.», kann der Autor bekräftigen, dass es leider vielerorts genau so aussieht. Ein zentral gepflegtes Data Warehouse, das sich seit Jahren nicht mehr wesentlich weiterentwickelt, weil die zentrale IT keine Ressourcen hat. Ein Schatten-Data Lake, der entstanden ist, weil man nicht mehr länger auf das Data Warehouse warten möchte. Unterschiedliche BI-Tools im Einsatz. Eine Vielzahl von Excel-Files, mit teils dutzenden Tabellenblättern, welche mit viel manuellem Aufwand bearbeitet werden und geschäftskritische Daten enthalten. Datensilos und Teams, die nicht miteinander arbeiten. Abweichende Kennzahlenwerte zwischen Reports. Ein neu eingeführtes agiles Framework, das bürokratischer und ineffizienter ist als alle Organisationsformen zuvor. Und schliesslich eine unzureichende Datenqualität, unter anderem weil klare Verantwortlichkeiten fehlen.

In den folgenden Kapiteln werden wir uns auf analytische Daten konzentrieren und darstellen, wie entsprechende Datenflüsse- und bestände gesamtunternehmerisch technisch implementiert und organisiert werden können.

3.2 Data Warehouse, Data Lake, Lakehouse

Die Art und Weise, wie analytische Daten aufbereitet und gespeichert werden, hat sich in den letzten Jahrzehnten zusammen mit den technischen Möglichkeiten und deren Implementierungskosten stetig weiterentwickelt. Diese Entwicklung kann am besten anhand der drei populärsten Architekturen beschrieben werden: dem Data Warehouse, dem Data Lake und dem modernen Ansatz des Lakehouse.

3 Datenarchitektur

Bevor wir die jeweiligen Ansätze in der Folge genauer betrachten, muss Folgendes erwähnt werden: Ob überhaupt einer dieser Architekturen implementiert wird, ist heute aufgrund der Mächtigkeit eines BI-Tools, wie z.B. Power BI, nicht mehr zwingend notwendig – insbesondere für kleinere Unternehmen oder KMU. In der Regel werden diese Ansätze dann relevant, sobald der Umfang der Daten sowie die Komplexität und Anzahl der analytischen Anwendungen steigt. Möchte man also einfach nur einen isolierten, klassischen Kostenstellenbericht mit Daten aus dem ERP abbilden, braucht man sehr wahrscheinlich kein Lakehouse und schon gar nicht ein Data Warehouse.

Data Warehouse

Das Data Warehouse oder kurz DWH ist die älteste Form der Architektur für analytische Daten. Es werden regelmässig (meist täglich) Daten aus Quellsystemen extrahiert, transformiert und in das DWH geladen (ETL-Prozess). Es werden Stamm- von Faktentabellen unterschieden und die entsprechenden

Tabellenschemas²⁵ sind fest definiert. Es liegt ein besonderes Augenmerk auf einer möglichst hohen Datenqualität, und es werden ausschliesslich strukturierte Daten abgespeichert.

Typische Stamm- und Faktentabellen im Finanzumfeld sind: Kontenplan, Kostenstellenstruktur, Produktstamm, Kundenstamm, FIBU-Transaktionen, Vertriebsbelege, Einkaufsbelege etc.

Ein DWH wird in der Regel von einer zentralen IT-Stelle gepflegt. Änderungen sind meist nur mit längeren Wartezeiten möglich (Bottlenecks). Es sind also eher starre Konstrukte mit fixierten Regeln. Datenabfragen erfolgen meist mit SQL oder direkt mit dem BI-Tool, welches letztendlich aber auch SQL-Abfragen an die Datenbank sendet. Während in der Vergangenheit ein DWH meist On-Premise implementiert wurde, setzen sich heute zunehmend Cloud-Lösungen durch. Die Datenspeicherung und Datenberechnung (Abfragen bzw. «Compute») finden an der gleichen Stelle auf dem DWH-Server statt.

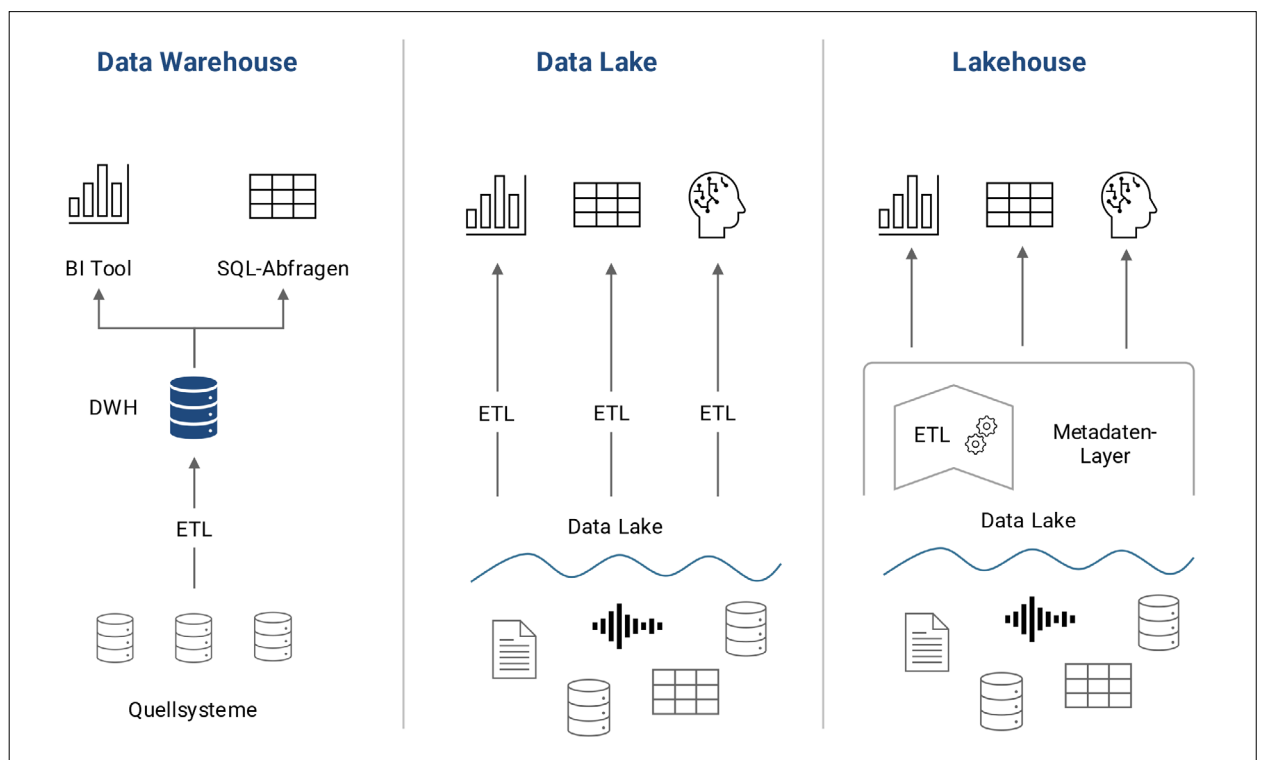


Abbildung 25: Übersicht der wichtigsten Architekturen (eigene Darstellung)

²⁵ Die Definition der Tabellenspalten: Spaltenbezeichnung, deren Inhalt und Datentyp (z.B. Text, ganze Zahl, Wahr/Falsch etc.)

3 Datenarchitektur

In der Tendenz sind die Implementation und der Unterhalt eines klassischen DWHs ein eher teures Unterfangen, womit kleine oder mittelgrosse Unternehmen in der Vergangenheit von diesem Vorgehen oft ausgeschlossen waren.

Data Lake

Mit der zunehmenden Popularität von unstrukturierten Daten und schnell skalierbaren Cloud-Speichertechnologien wurde der Ansatz des Data Lakes relevant. Grundsätzlich ist der Data Lake ein Speichertopf für alle Arten von strukturierten und unstrukturierten Daten, welche für die verschiedensten Anwendungen abgefragt und kombiniert werden können.

Der Data Lake ist in einigen Punkten das Gegenteil des DWH: Es gibt keine vordefinierten Tabellenschemas, das heisst, das Tabellenschema entsteht erst während der jeweiligen Abfrage der Daten. Weil es keine solchen Regeln und damit auch keinen einheitlichen ETL-Prozess gibt, ist die Implementierung eines Data Lakes im Vergleich eher günstig. Es entstehen aber implizite Kosten, verursacht durch die Unstrukturiertheit, welche mit einem Data Lake oftmals daherkommt.

Hinsichtlich Abfragesprachen ist man nicht mehr limitiert auf SQL, sondern kann z.B. auch mit Python direkt arbeiten, was gerade im Machine-Learning-Bereich die Norm ist. Auch ein wichtiges Unterscheidungsmerkmal ist die Trennung von Speicherung und Berechnung («Compute»), welche unabhängig voneinander skaliert werden können.²⁶

Lakehouse

Das DWH ist zwar zuverlässig, aber starr und teuer. Der Data Lake ist günstig und flexibel, aber meist unstrukturiert und unübersichtlich. Hier soll der heute moderne Ansatz des Lakehouse helfen und die Vorteile der jeweiligen Ansätze kombinieren.

Mit dem Lakehouse werden die Daten im meist günstigen und skalierbaren Cloudspeicher gespeichert (gleiche Technologie wie Data Lake). Gleichzeitig werden aber Tabellenschemas in einem Metadaten-Layer definiert, sodass die Abfrage anhand strukturierter Tabellen erfolgen kann. Diese Metaschicht ist flexibel und schnell erweiterbar, beschreibt die zugrundeliegenden Daten und es kann mit klaren Verantwortlichkeiten für die jeweiligen Logiken gearbeitet werden.

Hinsichtlich Abfragesprachen ist man frei wie beim Data Lake (SQL, Python etc.). Weiterhin wird meist die Technologie der sogenannten *Delta-Tables*²⁷ angewandt, womit man einen historischen Verlauf der jeweiligen Datenbestände automatisch abbildet.

Ein Lakehouse wird heute mit einer modernen Datenplattform implementiert, welche bereits eine Vielzahl an flexiblen Datenaufbereitungstools (Code oder Low-Code) mit sich bringt.

Die technologische Entwicklung hin zum Lakehouse ist sehr attraktiv für kleine und mittelgrosse Unternehmen, denn nun kann ohne riesige Projektbudgets von den Vorteilen eines Data Warehousing profitiert werden.

²⁶Man kann z.B. eine relativ einfache Abfrage auf einem grossen Datenbestand ausführen (z.B. eine Summe bilden) oder eine sehr komplexe, mehrstufige Abfrage auf einem kleinen Datenbestand (z.B. ein KI-Modell trainieren).

²⁷Vgl. Databricks (2025).

3 Datenarchitektur

	Data Warehouse	Data Lake	Lakehouse
Grundprinzip	Strukturiert, vertrauenswürdig	Alle Arten von Daten, Flexibel	Flexibel und vertrauenswürdig
Datenspeicherung	Physisch in der Datentabelle	Fileablage (sog. Blob-Storage)	Fileablage (sog. Parquet-Files)
Datenabfrage	SQL	Frei (SQL, Python etc.)	Weitestgehend frei (wichtigste Sprachen sind unterstützt)
Datentypen	Nur strukturierte Tabellen	Strukturiert, unstrukturiert	Strukturiert über Metaschicht
Tabellenschema	Beim Einspeisen der Daten	Beim Auslesen der Daten	Beides möglich
Anwendung	Standardberichte, in vielen Unternehmen heute ein «Legacy-System»	Ad-Hoc Analysen, Machine Learning / KI	BI, Ad-Hoc, Machine Learning / KI etc.
Vorteile	Vertrauenswürdig, Standardisiert	Kostengünstig, flexibel	Flexibel, skalierbar, vertrauenswürdig, meist eine integrierte Plattform
Nachteile	Starr, teuer, nur SQL	Ungeordnet, Abfragen und Tabellenschemata werden immer wieder neu definiert, keine Standardisierung	Braucht i. d. R. Cloudplattform

Abbildung 26: Vergleich DWH, Data Lake und Lakehouse

Welche Architektur ist die richtige für das eigene Unternehmen und wie wird sie implementiert? Wir werden in Kapitel 3.4 sehen, dass sich dies mit den heute modernen technologischen Möglichkeiten fast von selbst ergibt, denn das Lakehouse setzt sich aufgrund von dessen Vorteilen zunehmend

durch. Bevor dies jedoch genauer angeschaut wird, müssen zwei sehr wichtige Merkmale einer modernen Datenarchitektur betrachtet werden: Self-Service und Low-Code Tools.

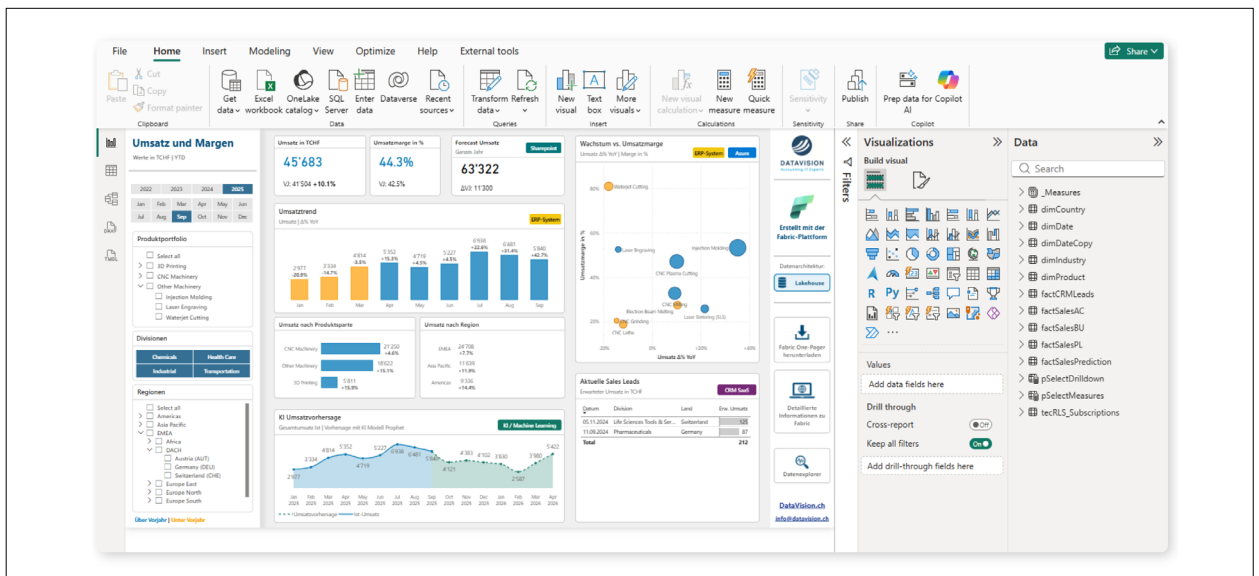


Abbildung 27: Low-Code Benutzerinterface von Power BI (eigene Darstellung)

3 Datenarchitektur

3.3 Die hohe Relevanz von Self-Service und Low-Code Tools

BI-Tools und generell Datentools haben sich im letzten Jahrzehnt rasant entwickelt. Heute ist es realistisch, eine umfangreiche BI-Lösung zu implementieren, und dabei nur minimal oder sogar gar keinen Code zu schreiben. Für eine moderne Datenplattform gehört es zum Standard, dass man alleine

über das User Interface (UI) mit Buttons, Dropdowns, Drag-and-Drop sowie einer meist einfach nutzbaren Formelsprache eine spezifische Lösung bauen kann. Diesen Ansatz nennt man Low-Code.

Und wenn es doch erforderlich wird, können die marktführenden Tools dennoch auch Code verarbeiten, z.B. SQL oder Python. Wenn man dabei einer Best Practice folgt, ergibt das in Summe eine hervorragende integrierte Lösung, welche man jederzeit weiter ausbauen kann.

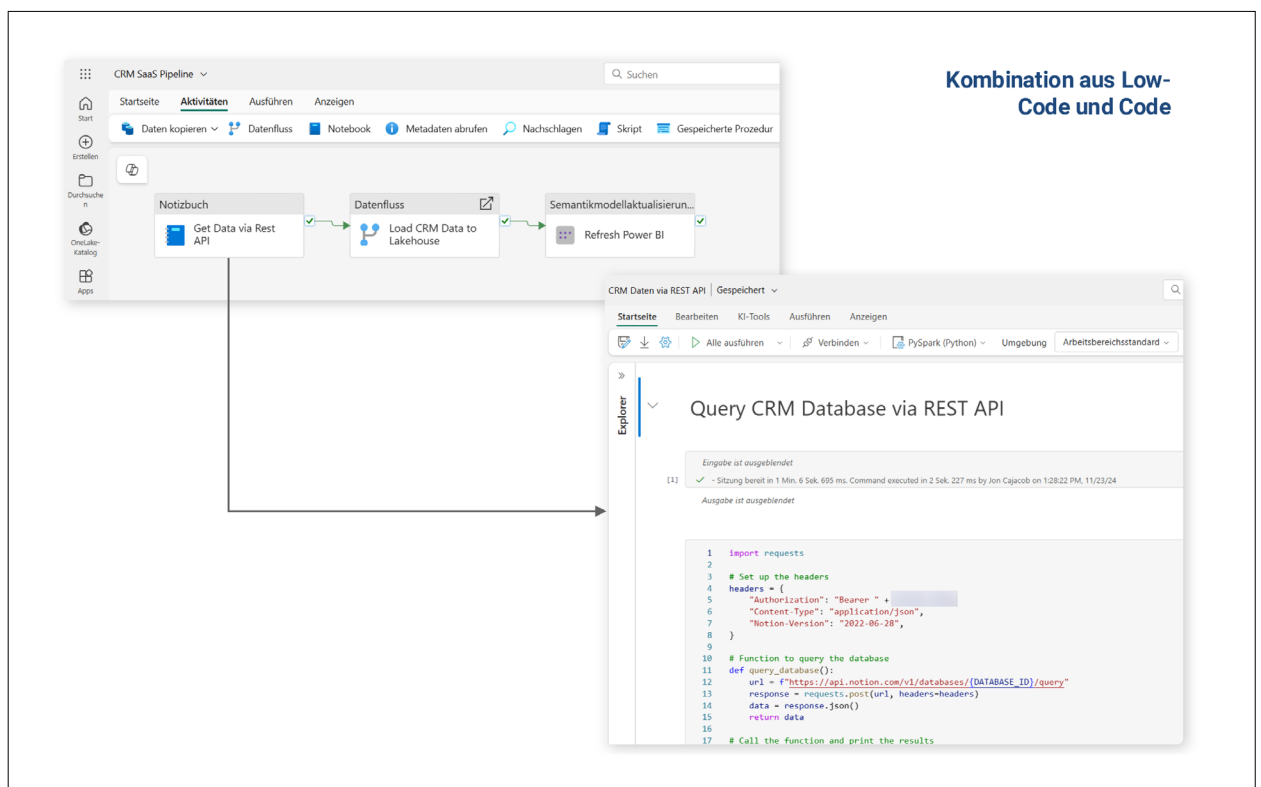


Abbildung 28: Kombination von Low-Code und Code mit Fabric: Ein Python Notebook (Code) wird per Pipeline (Low-Code) ausgeführt als Teil eines automatisierten Prozesses (eigene Darstellung)

3 Datenarchitektur

Beispiel SQL-Abfrage vs. Power Query

Um den Unterschied zwischen Code und Low-Code besser zu veranschaulichen, wird ein einfaches Beispiel anhand von klassischen Umsatzdaten betrachtet. Die folgende Tabelle speichert Umsatz- und Umsatzkostendaten nach Datum, Land, Kundenbranche und Produkt.

Wir wollen den aggregierten Umsatz sowie den Bruttogewinn nach Länderregion und Jahr abfragen, sortiert zuerst nach Jahr und dann Länderregion. Die Länderregion müssen wir aus einer separaten Stammdatentabelle beziehen, in welcher jedes Land einer Region zugeordnet ist.

	Date	ABC CountryID	ABC IndustryID	ABC ProductID	12F SalesAmount	12F CostAmount
1	2024-01-31	DE	Industry01	MG01	7280.3266	-4723.2282
2	2024-01-31	DE	Industry02	MG01	5244.282	-3402.3118
3	2024-01-31	DE	Industry03	MG01	4459.9052	-2893.4347
4	2024-01-31	DE	Industry04	MG01	2331.8932	-1512.853
5	2024-01-31	DE	Industry05	MG01	1101.678	-714.7312
6	2024-01-31	DE	Industry06	MG01	3493.0474	-2266.1703
7	2024-02-29	DE	Industry01	MG01	5639.6896	-3658.8388
8	2024-02-29	DE	Industry02	MG01	4062.472	-2635.5936
9	2024-02-29	DE	Industry03	MG01	3454.8561	-2241.3931
10	2024-02-29	DE	Industry04	MG01	1806.3962	-1171.9284
11	2024-02-29	DE	Industry05	MG01	853.4125	-553.665
12	2024-02-29	DE	Industry06	MG01	2705.8818	-1755.484
13	2024-03-31	DE	Industry01	MG01	9331.1228	-6053.715
14	2024-03-31	DE	Industry02	MG01	6721.5446	-4360.7094
15	2024-03-31	DE	Industry03	MG01	5716.2165	-3708.4868
16	2024-03-31	DE	Industry04	MG01	2988.7646	-1939.0088
17	2024-03-31	DE	Industry05	MG01	1412.0098	-916.0639
18	2024-03-31	DE	Industry06	MG01	4477.0044	-2904.5281
19	2024-04-30	DE	Industry01	MG01	8510.8044	-5521.5203
20	2024-04-30	DE	Industry02	MG01	6130.6396	-3977.3504
21	2024-04-30	DE	Industry03	MG01	5213.692	-3382.466
22	2024-04-30	DE	Industry04	MG01	2726.016	-1768.5465
23	2024-04-30	DE	Industry05	MG01	1287.8771	-835.5308
24	2024-04-30	DE	Industry06	MG01	1093.4316	-754.185

Abbildung 29: Beispiel Umsatzdaten (eigene Darstellung)

Zuerst die Abfrage mit SQL:

```
1 SELECT
2     FORMAT(sales.[Date], 'yyyy') AS [Year],
3     country.[Region],
4     SUM(sales.[SalesAmount]) AS [TotalSales],
5     SUM(sales.[SalesAmount] + sales.[CostAmount]) AS [TotalMargin]
6 FROM
7     [dbo].[factSales] sales
8 INNER JOIN
9     [dbo].[dimCountry] country
10    ON sales.[CountryID] = country.[CountryID]
11 GROUP BY
12     FORMAT(sales.[Date], 'yyyy'),
13     country.[Region]
14 ORDER BY
15     [Year] DESC,
16     country.[Region]
```

Abbildung 30: SQL-Abfrage (eigene Darstellung)

3 Datenarchitektur

Zuerst werden die benötigten Spalten selektiert (SELECT) und dann holen wir die Regionen-Information aus der Ländertabelle (INNER JOIN). Anschlies-

send aggregieren wir die Werte nach Jahr und Region (GROUP BY) und als letzter Schritt sortieren wir die Ergebnistabelle (ORDER BY).

Resultat:

	ABC Year	ABC Region	ABC TotalSales	ABC TotalMargin
1	2025	Americas	12'089'865	5'177'750
2	2025	Asia Pacific	15'086'888	6'533'515
3	2025	EMEA	32'228'152	14'699'104
4	2024	Americas	10'159'791	3'993'063
5	2024	Asia Pacific	12'892'302	5'234'209
6	2024	EMEA	28'970'450	12'748'262
7	2023	Americas	9'031'959	3'381'658
8	2023	Asia Pacific	11'626'949	4'536'474
9	2023	EMEA	26'977'070	11'550'708
10	2022	Americas	8'507'627	3'092'515
11	2022	Asia Pacific	11'006'332	4'176'774
12	2022	EMEA	26'377'060	11'233'791

Abbildung 31: Ergebnis der Abfrage mit SQL (eigene Darstellung)

Und das gleiche Ergebnis können wir mit Power Query in Power BI oder Excel erzielen:

The screenshot shows the Power Query Editor interface. The main area displays a table with the following data:

	Year	Region	TotalSales	TotalMargin
1	2025	Americas	12'089'864.89	5'177'749.68
2	2025	Asia Pacific	15'086'887.87	6'533'514.59
3	2025	EMEA	32'228'152.27	14'699'104.47
4	2024	Americas	10'159'790.51	3'993'062.76
5	2024	Asia Pacific	12'892'301.69	5'234'209.32
6	2024	EMEA	28'970'450.02	12'748'261.66
7	2023	Americas	9'031'958.81	3'381'657.74
8	2023	Asia Pacific	11'626'948.52	4'536'474.22
9	2023	EMEA	26'977'070.00	11'550'707.99
10	2022	Americas	8'507'626.88	3'092'515.08
11	2022	Asia Pacific	11'006'331.94	4'176'774.38
12	2022	EMEA	26'377'060.35	11'233'791.28

The interface includes a ribbon with tabs: Date, Home, Transform, Add Column, View, Tools, and Help. The Home tab is active, showing options like Close & Apply, New Source, Recent Sources, Enter Data, Data source settings, Manage Parameters, Export query results, Refresh Preview, Advanced Editor, Choose Columns, Remove Columns, Reduce Rows, Sort, Split Column, Group By, Data Type: Text, Use First Row as Headers, and Combine. The right-hand pane shows the Query Settings for 'factSalesAC', including the Name, All Properties, and a list of Applied Steps: Source, Navigation, Removed Other Columns, Added Custom, Inserted Year, Merged Queries, Expanded dimCountry, Grouped Rows, Changed Type, and Sorted Rows.

Abbildung 32: Ergebnis der Abfrage mit Power Query in Power BI

3 Datenarchitektur

In der Abbildung 32 sieht man auf der rechten Seite die Transformationsschritte, welche nach und nach auf die Daten angewandt wurde. Diese werden automatisch erstellt, sobald über das UI eine Transformation vorgenommen wird. Im Prinzip hat man einen Automatisierungsprozess erstellt, welcher beim Laden (oder Aktualisieren) der Daten im BI jeweils angewandt wird (ETL-Prozess).

Das Beispiel zeigt: Beide Wege führen zum gleichen Ziel. Es liegt dabei auf der Hand, dass dank des Low-Code-Ansatzes mächtige Datentools von IT-affinen Personen in der Praxis angewandt werden können.

Die Konsequenz von Low-Code für eine moderne Datenarchitektur

Damit spannt sich der Bogen zum Thema Datenarchitektur. Low-Code hat deshalb eine so hohe Relevanz, weil es den sogenannten Self-Service ermöglicht: Fachbereiche wie z.B. das Controlling oder die Fertigung können eigenständig Datenprozesse und interaktive Dashboards entwickeln, und dabei nur minimal abhängig von einer zentralen IT-Stelle sein. Dies erhöht die Eigenständigkeit, Geschwindigkeit und Flexibilität für die jeweiligen Teams. Voraussetzung ist aber, dass genügend IT-Affinität und Kapazität im entsprechenden Team vorhanden ist.

Und so kann allgemein ein Trend hin zur dezentralisierten Datenarchitektur beobachtet werden. Die zentrale IT stellt dabei die Plattform und Infrastruktur bereit und verantwortet übergreifende Themen wie Security, Lizenzen und Skalierbarkeit. Auf dieser Grundlage erstellen die individuellen Teams bzw. Domänen ihre Datenprodukte selbstständig.

Dieses Prinzip der Dezentralisierung wird auch Data Mesh genannt. Die Architektur wird in sogenannte Daten-Domänen organisiert (z.B. Finance, HR, Supply Chain Management etc.), welche eigenständig Datenprodukte (BI-Dashboards, KI-Modelle, RPA-Prozesse etc.) entwickeln und betreiben.²⁸

Es gilt zu erwähnen, dass es auch hier keine Schwarz-Weiss-Trennung gibt zwischen zentraler IT-Stelle und dezentralem Fachbereich. Es gibt beispielsweise Szenarien, wo die zentrale Stelle bereits eine Vorselektion an Datenquellen auf der Plattform zur Verfügung stellt, und die jeweiligen Teams können dann auf Basis dieser leicht vorbereiteten Rohdaten weiterarbeiten. Hierbei setzt sich zunehmend das Prinzip der Medallion-Architektur durch, bei der Daten über mehrere Stufen – Bronze, Silber und Gold – zunehmend veredelt und auch zwischengespeichert werden.²⁹

3.4 Umsetzung einer modernen Datenarchitektur mit einer Plattformlösung

Wir wissen nun, wieso eine solide Datenarchitektur für analytische Anwendungen wichtig ist, und kennen die verschiedenen Ansätze. Wie aber geht man vor, um eine moderne Architektur für die eigene Organisation zu implementieren?

Wichtig hierbei ist die Entscheidung für eine umfangreiche und flexible Plattformlösung. Ziel ist eine gemeinsame Arbeitsumgebung für alle Mitwirkenden. Ansonsten hat man eine schwer kontrollierbare Landschaft mit auf dem lokalen PC ausgeführten Python-Notebooks, On-Premise gehosteten SQL-Datenbanken und in der Cloud publizierten Power BI-Berichten.

Einen Startpunkt definieren

In der Realität hat man selten eine grüne Wiese als Ausgangslage, sondern vielleicht bereits schon ein BI-Tool mit diversen Berichten im Einsatz oder sogar noch ein älteres Data Warehouse auf einem lokalen Server installiert. Es ist zu empfehlen, zuerst die Ausgangslage genau zu verstehen, das heisst am besten diese visuell aufzuzeichnen, sodass die bestehenden Systeme, Datenbestände und besonders auch Datenflüsse klar sind.

²⁸ Vgl. Dehghani (2022), S. 35 ff.

²⁹ Vgl. Databricks (2025).

3 Datenarchitektur

Um dann einen konkreten Startpunkt zu finden, lohnt es sich, Use-Cases zu definieren und zu priorisieren und dann mit einem solchen zu beginnen. Entsprechende Klassiker im Finanzumfeld sind der Kostenstellenreport oder der Sales-Report. Diese bringen bereits einen grossen Mehrwert und sind von der Komplexität her meist in einem vernünftigen Mittelfeld.

Datenplattform evaluieren

Heute ist es eher die Norm als die Ausnahme, dass eine moderne Datenplattform alle relevanten Funktionen aus dem Bereich Data & Analytics abdeckt. Denn es ist nicht empfehlenswert, für verschiedene Themen, z.B. BI oder Machine Learning, mehrere individuelle Softwarelösungen anzuschaffen, da man dadurch u.a. Schnittstellenproblematiken schafft.

Eine umfangreiche Datenplattform deckt mindestens die folgenden Funktionen ab:

- Grosse Anzahl von Datenkonnektoren für allerlei denkbare Quellsysteme
- Mächtige und flexible Datenintegrationstools (Datentransformation, ETL)
- Data Warehousing (klassisch SQL), Data Lake und Lakehouse
- Business Intelligence (Modellierung und Visualisierung)

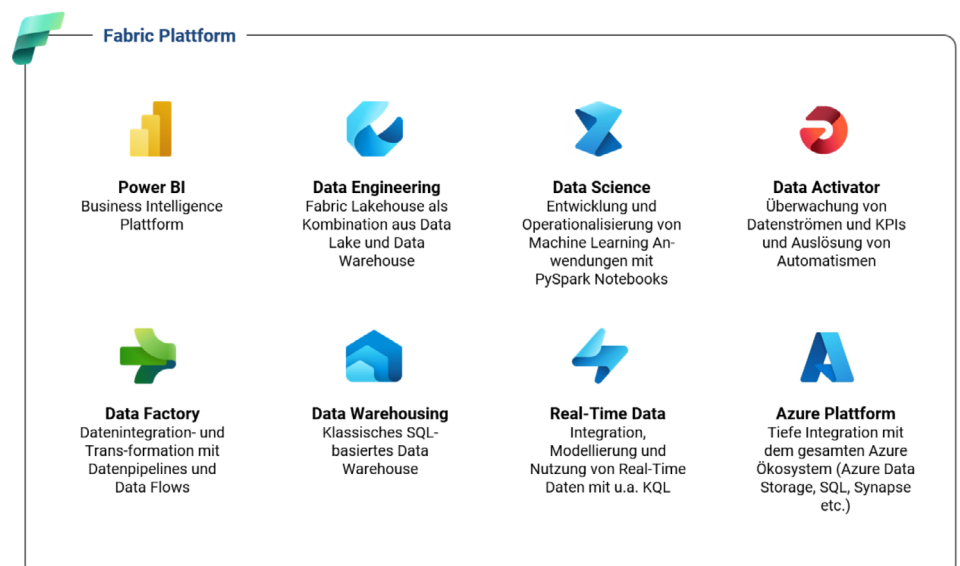
- Möglichkeit, Python-Notebooks zu erstellen und zu operationalisieren (und damit Machine-Learning-Anwendungen zu implementieren)
- Integration von Real-Time-Datenströmen
- Umfangreiche Governance-Funktionen (Berechtigungen, Datenklassifizierung, Datenkatalogisierung etc.)

Und all diese Punkte sollten sowohl ein Low-Code wie aber auch Code-Vorgehen unterstützen.

Ausserdem ist es generell zu empfehlen, mit einer einfach skalierbaren Cloudlösung zu arbeiten. Wobei angemerkt werden muss, dass eine lokale Installation einer Plattform mit dem oben genannten Funktionsumfang heutzutage kaum mehr realistisch ist.

Um ein konkretes Beispiel zu nennen, zeigt die folgende Darstellung die Module der Datenplattform Fabric von Microsoft. Power BI ist dabei nur ein Teil von vielen Komponenten. Alle Funktionen finden sich dabei direkt im Browser, das heisst, man muss nicht zwischen verschiedenen Applikationen hin- und herwechseln. Die Module sind sehr gut miteinander integriert, das heisst, Übergänge von Daten von z.B. einem Lakehouse in das Power BI sind nahtlos möglich.

Abbildung 33:
Übersicht der Module der Datenplattform Fabric (eigene Darstellung)



3 Datenarchitektur

Proof-of-Concept (PoC) durchführen

Hat man einen Use-Case gefunden, soll dieser im Rahmen eines Proof-of-Concept (PoC) auf der neuen Datenplattform umgesetzt werden. Dabei ist eine wichtige Arbeitsannahme, dass die Lösung zu Beginn noch nicht in allen Aspekten ausgereift sein wird. Auch hier empfiehlt es sich, ein iteratives (agiles) Vorgehen zu wählen.

Ist der PoC umgesetzt, muss entschieden werden, ob mit der gewählten Plattform weitergearbeitet werden soll.

Einführung und Operationalisierung der neuen Datenplattform

Hat sich der PoC bewährt und möchte man die Plattform für eine moderne Datenarchitektur einführen, müssen einige Entscheidungen getroffen werden, u.a.:

- Wo sollen unsere Daten gehostet werden, sofern wir uns für eine Cloudlösung entscheiden? I. d. R. ist dies der geografische Standort (Region) des Unternehmens.
- Wie werden die Arbeitsbereiche der Teams strukturiert? Bekommt jedes Team einen eigenen Wirkungsbereich oder wird alles an zentraler Stelle konsolidiert? Wie in früheren Abschnitten erwähnt, hat der erste Ansatz viele Vorteile für Organisationen, die wirklich datengetrieben arbeiten wollen.
- Wird das klassische, SQL-basierte Data Warehouse benutzt oder ein Lakehouse? Soll eine Medallion-Architektur mit Bronze, Silber und Goldstufen angewandt werden? Gibt es ein zentrales Lakehouse oder mehrere dezentrale? Der Trend bei modernen Datenplattformen zeigt, dass sich das Lakehouse zunehmend durchsetzt, entsprechend liegt der Fokus bei der Weiterentwicklung und dem Funktionsumfang der Plattform. Hinsichtlich Medallion-Architektur empfiehlt es sich, diverse Stamm- und Faktentabellen möglichst unternehmensweit zu standardisieren, sodass nicht jedes Mal z.B. eine Kundenstammtabellen von neuem erstellt werden muss.

- Wer ist verantwortlich für den fehlerfreien Betrieb eines bestimmten Datenprodukts? Wer leistet Support, sollte eine Datenaktualisierung fehlschlagen? Auch hier ist entscheidend, ob man einen eher dezentralen oder zentralen Ansatz fährt.
- Wie werden die Berechtigungen gesteuert? Wer ist verantwortlich für Freigaben auf Datenbeständen und Berichten? Auch hier gibt es zentrale und dezentrale Ansätze oder ein Mix aus beidem.
- Wie werden Daten und generell Datenobjekte (wie z.B. Dashboards, Modelle etc.) strukturiert auffindbar gemacht? Gute Plattformlösungen bieten verschiedene Funktionen für das Katalogisieren und Beschreiben von Datenprodukten an.
- Welche Funktionen sind generell erlaubt, welche nur für bestimmte Usergruppen und welche gar nicht? Das Herunterladen von Daten aus Dashboards in Excel beispielsweise könnte nur für bestimmte Personen erlaubt sein. Das Nutzen von Drittdiensten, z.B. «Bing Maps», könnte ganz deaktiviert werden.

Kontinuierlich verbessern

Wie bereits erwähnt, wird es nicht möglich sein, all die eben genannten Fragen am Anfang endgültig beantworten zu können. Auch können sich Rahmenbedingungen und damit Anforderungen mit der Zeit verändern. Der Vorteil einer (guten) Datenplattform ist deren Skalierbar- und Anpassbarkeit. Deshalb: agil vorgehen, der 80:20-Regel folgen und das Setup kontinuierlich verbessern.

3 Datenarchitektur

Case Study E-Velo AG: Datenarchitektur für eine datengetriebene Zukunft

Die E-Velo AG mit Sitz in Basel produziert und vertreibt E-Bikes von hohem Qualitätsstandard in der Schweiz. Die Bikes werden auf Auftrag gefertigt, und es gibt nur einen sehr kleinen Lagerbestand an vorgefertigten Produkten. Die Firma ist grundsätzlich profitabel, es ist aber nicht genau klar warum, bzw. welcher Teil des Produktportfolios positiv und welcher negativ dazu beiträgt. Das Controlling exportiert einmal im Monat manuell Buchungsdaten per CSV aus dem ERP, welches dann in einem Excel-Template mit einem VBA-Makro aufbereitet wird. Das Ergebnis, eine Liste mit fast 1000 Zeilen, zeigt Kontensalden nach Kostenstelle, und wird dann an die Geschäftsleitung per E-Mail verschickt. Das Makro wurde von einer Person erstellt, welche heute nicht mehr im Unternehmen tätig ist. Neben dem ERP nutzt die Firma ein Cloud-CRM-Tool für die Erfassung der Kundenstammdaten. Diese CRM-Daten sind damit extern gespeichert.

Die Geschäftsleitung hat die Initiative gestartet, das Unternehmen datengetriebener auszurichten. Ein erster Use-Case ist ein automatisierter Bericht, welcher die Profitabilität des Produktportfolios darstellt. Die dafür notwendigen Daten sollen dann später die Basis bilden für weiterführende Anwendungsfälle, wie z. B. einen KI-unterstützten Umsatzforecast.

Um diesen ersten Use-Case zu implementieren, wird zuerst die Datenbasis geschaffen und verbessert: Neu werden im ERP-System Arbeitszeiten aus der Produktion und Materialeinsätze einem Projekt bzw. Auftrag direkt zugeordnet. So wird es möglich, einen Deckungsbeitrag pro Produkt recht genau herzuleiten.

Es wurde entschieden, die Datenplattform Fabric einzuführen. Für den Setup der Plattform ist die zentrale IT zuständig, welche zusätzlich externen Support einholt, sodass von Anfang an mit Best Practices gearbeitet wird. Die zentrale Stelle erstellt ein Bronze-Lakehouse und lädt eine Vorselektion der notwendigen Quelldaten aus dem ERP sowie aus dem CRM-System. Ersteres wird per SQL-Query abgefragt, letzteres per Rest API durch ein Python-Notebook. Im Controlling-Team gibt es eine IT-affine Person, welche Erfahrung hat mit Low-Code-Tools. Es wurde definiert, dass fortan ein wesentlicher Teil des Pensums dieser Person für Daten- und Automatisierungsthemen reserviert wird. Sie wird durch einen externen Experten gecoacht. Diese Person erstellt ein Gold-Lakehouse im Controlling-Arbeitsbereich und erstellt die notwendigen Dimensions- und Faktentabellen: Kundenstamm, Produktstamm, Kostenstellenstamm, Vertriebsbelege und Umsatzkostenbelege.

Darauf aufbauend wird ein Power BI-Semantikmodell erstellt, welches diese Tabellen miteinander in Beziehung setzt und die Basis für die Berechnungen bzw. KPIs (Measures) bildet. Zu guter Letzt werden die Berichte wie von der Geschäftsleitung gewünscht erstellt und publiziert.

Die Datenaktualisierung wird gesamthaft mit einer Datenpipeline in der Nacht ausgeführt. Dabei werden die notwendigen Daten aus dem ERP extrahiert sowie die CRM-Daten aus der externen SaaS-Applikation. Die Daten werden über zwei Stufen (Bronze und Gold) veredelt bis zur Verwendung im Dashboard. Zukünftig bilden die Gold-Tabellen auch die Basis für weitere Anwendungsfälle.

Die Datenqualität der wichtigsten Datenfelder wird in der Zwischenzeit mit einer weiteren Berichtsseite dargestellt und fortan laufend beurteilt und verbessert.

3.5 Datenqualität: Strategien zur kontinuierlichen Verbesserung

Das Thema Datenqualität ist von hoher Relevanz und leider gibt es kein Geheimrezept, dies schnell und einfach zu «erledigen». Vielfach wird Datenqualität im Kontext der Gesamtarchitektur betrachtet, weshalb sich dieses kurze Kapitel mit ausgewählten Strategien zur Bewältigung der Problematik befasst.

Ursachen für ungenügende Datenqualität

Wir wollen hier zwei wichtige Gründe für eine ungenügende Datenqualität hervorheben: Erstens verändern sich Daten ständig. Es werden laufend neue Daten erfasst und bestehende Daten verändert, entweder durch den Menschen oder maschinell. Dies macht Datenqualität zum Dauerthema, das man nicht einmalig erledigt und dann beiseitelegt. Zweitens ist das Thema Datenqualität nicht besonders attraktiv, wie etwa im Vergleich zu einem vielversprechenden KI-Projekt oder einem schön gestalteten Dashboard. Das hat zur Folge, dass sich oftmals niemand so wirklich intensiv damit befassen möchte, bzw. das Thema bekommt selten die notwendige Aufmerksamkeit.

Wieso es sich lohnt, an der Datenqualität zu arbeiten: Schlechte Datenqualität heisst: Daten sind (zu) ungenau, inkonsistent, veraltet oder schlicht komplett falsch. Und das kann für eine Organisation verheerende Konsequenzen haben. Man stelle sich vor, einem Produkt werden zu viele Kosten zugeordnet, weil eine Mengenangabe auf der Stückliste falsch ist, und dann wird entschieden, das Produkt aus dem Portfolio zu entfernen, obwohl es eigentlich profitabel wäre.

Es ist wichtig zu verstehen, dass Datenprodukte aller Art nur so gut sind, wie die Qualität der Daten, mit denen wir diese befüllen («garbage-in, garbage-out»). Dies gilt für Berichte, Dashboards und insbesondere auch KI-Projekte, wo Vorhersagen gemacht werden.

Ansätze zur kontinuierlichen Verbesserung

Die folgenden Massnahmen können helfen, die Datenqualität laufend zu verbessern oder zumindest ein bestimmtes Niveau beizubehalten. Es empfiehlt sich, die Massnahmen zu kombinieren.

- Datenqualität messen und transparent machen: Ein vorhandenes BI-Tool kann genutzt werden, um die Datenqualität mittels eines Dashboards transparent darzustellen. Das Messen, mithilfe von Kennzahlen, ist nicht immer einfach; es gibt aber praktikable Ansätze, eine Quantifizierung vorzunehmen (z.B. % der E-Mail-Adressen im CRM, welche leer sind). Zugleich hat das BI selbst einen nützlichen Nebeneffekt zur Verbesserung der Datenqualität, weil im BI werden schon länger verborgene Probleme plötzlich sehr gut sichtbar, womit Handlungsdruck entsteht.
- Die wichtigsten und kritischsten Probleme identifizieren und gezielt priorisieren: In der Regel sind Ressourcen für die Verbesserung der Datenqualität begrenzt, sodass es empfehlenswert ist, sich auf diejenigen Probleme zu fokussieren, die die grössten negativen Auswirkungen auf Prozesse, Auswertungen oder Entscheidungen haben. Dabei kann die 80:20-Regel als Hilfestellung dienen: Oftmals verursacht ein vergleichsweise kleiner Teil der Probleme einen erheblichen Teil der Beeinträchtigungen.
- Ursachen für Probleme möglichst genau verstehen und an der Quelle beheben: Warum sind denn ein wesentlicher Teil unserer Kunden keiner oder der falschen Branche zugeordnet? Ist evtl. das entsprechende Dropdown-Menü schlecht bedienbar und passen irgendwie keine der Branchen so richtig? Es lohnt sich, den Ursachen für die kritischen Probleme genau auf den Grund zu gehen, weil so Probleme schon bei der Datenerfassung vorgebeugt werden können.
- Klare Verantwortlichkeiten für die Pflege und Qualität bestimmter Daten festlegen: Für zentrale Datenbestände sollten jeweils konkrete

3 Datenarchitektur

Personen oder Funktionen benannt werden, die für deren Pflege, Vollständigkeit, Aktualität und Qualität verantwortlich sind. So trägt beispielsweise die Buchhaltung die Verantwortung für eine präzise und aussagekräftige Erfassung von Geschäftstransaktionen, der Vertrieb für ein möglichst vollständiges und verlässliches CRM und HR für die Korrektheit und Aktualität der Mitarbeiterstammdaten, etwa in Bezug auf Ein- und Austrittsdaten oder Teamzugehörigkeiten. Der jeweilige Status der Datenqualität kann anschliessend über ein entsprechendes Dashboard transparent gemacht und nachvollzogen werden. Dieser Ansatz entspricht dem Grundgedanken von Data Ownership.³⁰

Keine dieser Massnahmen ist für sich genommen ein Patentrezept oder sollte isoliert angewendet werden. Auch wird die Verbesserung der Datenqualität immer mit erheblichem Aufwand verbunden bleiben. Entsprechende Bemühungen werden in einer zunehmend datengetriebenen Welt jedoch immer wichtiger und immer deutlicher belohnt.

³⁰ Vgl. DAMA International (2017), S. 68 ff.

Weitere Automatisierungsformen und Tools



Irfan Yasar

4.1 Prozesse und Automatisierungspotenzial

Nachdem sich diese Broschüre bisher mit der Automation von Datenaufbereitung und Reporting beschäftigt hat, werden nun weitere Formen der Automatisierung betrachtet. Der Fokus wird dabei erweitert von rein datengetriebenen Abläufen hin zu übergreifenden Geschäfts- und Arbeitsprozessen.

Der erste und entscheidende Schritt jeder Automatisierungsinitiative besteht darin, bestehende Arbeitsabläufe detailliert zu erfassen und Prozesse mit hohem Automatisierungspotenzial zu identifizieren. Nicht jeder Prozess eignet sich gleichermaßen für die Automatisierung – die systematische Analyse hilft, Prioritäten zu setzen und Quick Wins frühzeitig zu realisieren. Dabei geht es nicht darum, möglichst viele Prozesse zu automatisieren, sondern die richtigen Prozesse auszuwählen, bei denen der Nutzen den Aufwand deutlich übersteigt.

Kriterien für automatisierungsgerechte Prozesse

Die Eignung eines Prozesses für die Automatisierung lässt sich anhand mehrerer Kriterien beurteilen, die in ihrer Kombination ein Gesamtbild ergeben. Das wichtigste Kriterium ist die **Wiederholungsfrequenz**: Tätigkeiten, die regelmässig und in grosser Zahl anfallen, bieten das grösste Einsparpotenzial. Je häufiger ein Prozess durchgeführt wird, desto schneller amortisiert sich die Investition in seine Automatisierung. Ein Prozess, der täglich hundertfach ausgeführt wird, rechtfertigt einen deutlich höheren Implementierungsaufwand als einer, der nur monatlich anfällt.

Eng damit verbunden ist die Frage der **regelbasierten Logik**. Prozesse mit klaren «Wenn-Dann»-Regeln lassen sich gut in Algorithmen übersetzen. Wenn die Entscheidungslogik eindeutig dokumentiert werden kann – beispielsweise «Wenn der Rechnungsbetrag mit der Bestellung übereinstimmt und der Lieferant bekannt ist, dann zur Zahlung freigeben» – ist eine Automatisierung in der Regel technisch umsetzbar. Schwieriger wird es bei Prozessen, die implizites Wissen, Erfahrungswerte oder subjektive Einschätzungen erfordern.

Das **Datenvolumen** spielt ebenfalls eine zentrale Rolle. Aufgaben, die grosse Datenmengen verarbeiten, profitieren überproportional von Automatisierung, da manuelle Verarbeitung mit steigendem Volumen exponentiell aufwändiger wird. Während ein Mitarbeiter zehn Buchungen pro Stunde manuell prüfen kann, verarbeitet ein automatisiertes System die gleiche Menge in Sekunden – und skaliert problemlos auf tausende oder zehntausende Transaktionen.

Weitere begünstigende Faktoren sind **standardisierte Eingabe- und Ausgabeformate**, eine **geringe Ausnahmerequote** sowie **hohe Fehlerkosten**. Prozesse mit klar definierten Datenstrukturen sind leichter zu automatisieren als solche mit variablen oder unstrukturierten Informationen. Bei Prozessen mit vielen Sonderfällen ist oft eine Teilautomatisierung der pragmatischere Ansatz, bei dem Standardfälle automatisch und Ausnahmen manuell bearbeitet werden. Und schliesslich rechtfertigen Prozesse, bei denen manuelle Fehler zu signifikanten finanziellen oder regulatorischen Konsequenzen führen können, höhere Investitionen in Automatisierung – die Vermeidung eines einzigen kostspieligen Fehlers kann die gesamte Implementierung refinanzieren.

Im Finanz- und Rechnungswesen ergeben sich aus diesen Kriterien typische Automatisierungskandidaten: Die Kreditorenbuchhaltung mit Rechnungseingang, -prüfung und -freigabe weist alle

4 Weitere Automatisierungsformen und Tools

genannten Merkmale auf – hohe Frequenz, klare Regeln, standardisierte Formate. Gleiches gilt für die Debitorenbuchhaltung mit der Verarbeitung von Zahlungseingängen und dem Mahnwesen, für Kontenabstimmungen, die regelmässige Berichterstellung, Spesen-Abrechnungen, Intercompany-Abstimmungen etc.

Methodik zur Prozessanalyse

Eine systematische Prozessanalyse beginnt mit der Erstellung eines Prozessinventars, das alle relevanten Prozesse mit ihren wesentlichen Merkmalen erfasst:

- Wie häufig wird der Prozess durchgeführt?
- Wie lange dauert ein Durchlauf?
- Welche Systeme sind beteiligt?
- Wer ist verantwortlich?

Diese Übersicht ermöglicht eine erste Einschätzung, welche Prozesse näher betrachtet werden sollten.

Dann folgt eine **detaillierte Dokumentation**; dabei werden die einzelnen Arbeitsschritte, Entscheidungspunkte und Schnittstellen im Detail aufgenommen. Besonders wichtig ist es, auch informelle Arbeitsweisen und Workarounds zu erfassen – jene Abweichungen vom offiziellen Prozess, die es im Alltag gibt, und welche oft wertvolle Hinweise auf Schwachstellen oder fehlende Funktionalitäten geben.

Anhand der beschriebenen Kriterien wird anschliessend das **Automatisierungspotenzial bewertet**, typischerweise auf einer Skala von niedrig bis hoch. Dieser Bewertung wird eine **Aufwand-Nutzen-Analyse** gegenübergestellt: Dem ermittelten Nutzen in Form von Zeitersparnis, Qualitätsverbesserung und Fehlerreduktion werden die erwarteten Implementierungskosten und -risiken gegenübergestellt. Daraus ergibt sich eine **Priorisierung**, wobei es sich häufig empfiehlt, mit Prozessen zu beginnen, die hohes Potenzial bei moderatem Aufwand bieten – den sogenannten Quick Wins. Diese ermöglichen schnelle Erfolge und stärken die Akzeptanz für weitere Automatisierungsvorhaben.

Häufige Automatisierungshürden

Bei der Analyse werden regelmässig auch Hürden sichtbar, die einer Automatisierung im Wege stehen. Die häufigste Ursache für gescheiterte Automatisierungsprojekte ist **mangelnde Datenqualität**.

Unvollständige, inkonsistente oder fehlerhafte Daten führen dazu, dass automatisierte Systeme falsche Ergebnisse produzieren oder ständig Ausnahmen generieren, die manuell bearbeitet werden müssen. In vielen Fällen muss vor der eigentlichen Automatisierung zunächst ein Projekt zur Verbesserung der Datenqualität durchgeführt werden.

Eine weitere technische Hürde ist die **fehlende Systemintegration**. Wenn Prozesse über mehrere Systeme verteilt sind, die keine Schnittstellen zueinander haben, kann die Automatisierung technisch aufwändig werden. Moderne Integrationstechnologien wie APIs, Middleware-Lösungen oder RPA-Tools können hier Abhilfe schaffen, erfordern jedoch entsprechende Investitionen und Know-how.

Auf der fachlichen Seite stellen **unklare oder undokumentierte Prozesse** eine erhebliche Herausforderung dar. Implizites Wissen, das nur in den Köpfen erfahrener Mitarbeitender existiert, muss zunächst explizit gemacht und in formale Regeln übersetzt werden. Dieser Prozess der Wissensexternalisierung ist oft aufwändiger als die technische Implementierung selbst, führt aber zu einem besseren Verständnis der eigenen Prozesse und ermöglicht auch unabhängig von der Automatisierung Verbesserungen.

Schliesslich darf der **organisatorische Widerstand** nicht unterschätzt werden. Automatisierung verändert Arbeitsabläufe und Rollenbilder, was bei Mitarbeitenden Bedenken auslösen kann – etwa die Sorge vor Arbeitsplatzverlust oder Kontrollverlust über die eigenen Prozesse. Transparente Kommunikation über die Ziele der Automatisierung, die frühzeitige Einbindung der betroffenen Mitarbeitenden in Entscheidungsprozesse und gezielte Schulungsmassnahmen sind entscheidend, um Vorbehalte abzubauen und die Akzeptanz zu fördern.

4 Weitere Automatisierungsformen und Tools

Von der Analyse zur Umsetzung

Die Identifikation von Automatisierungspotenzial ist kein einmaliges Projekt, sondern sollte als kontinuierlicher Prozess verstanden werden. Mit jeder erfolgreich umgesetzten Automatisierung verändern sich die Rahmenbedingungen: Neue Schnittstellen entstehen, Mitarbeitende gewinnen Erfahrung mit automatisierten Prozessen, und zuvor als schwierig eingestufte Prozesse werden möglicherweise realisierbar. Zudem entwickeln sich die verfügbaren Technologien ständig weiter – was heute noch aufwändig erscheint, kann morgen durch neue Tools oder Plattformen deutlich einfacher umsetzbar sein.

Ein pragmatischer Ansatz besteht darin, mit einem **Pilotprojekt** zu beginnen, das überschaubar genug ist, um schnell Ergebnisse zu liefern, aber repräsentativ genug, um Erkenntnisse für zukünftige Projekte zu gewinnen. Die dabei gesammelten Erfahrungen – sowohl positive als auch negative – bilden die Grundlage für eine kontinuierliche Verbesserung der Automatisierungsstrategie.

Für kleine und mittlere Unternehmen empfiehlt sich dabei ein schrittweises Vorgehen:

Beginnen Sie mit der Analyse eines klar abgegrenzten Bereichs wie der Kreditorenbuchhaltung, identifizieren Sie zwei bis drei konkrete Prozesse mit hohem Potenzial, und setzen Sie diese als Pilotprojekte um. Die gewonnenen Erfolge und Erkenntnisse bilden dann die Basis für die Ausweitung auf weitere Bereiche des Unternehmens.

4.2 Robotic Process Automation (RPA)

Der Begriff **Robotic Process Automation** – kurz RPA – beschreibt die automatisierte Durchführung von informationstechnischen Tätigkeiten mit Softwareprogrammen, die ansonsten manuell von Menschen ausgeführt werden. Diese Softwareprogramme werden auch als Software-Roboter, «Bots» oder «Robotics» bezeichnet. Es handelt sich dabei ausdrücklich nicht um physisch existente Roboter,

wie sie für die Automatisierung direkter physischer Arbeitstätigkeiten in der Fertigung eingesetzt werden, sondern um rein digitale Entitäten, die auf Computersystemen laufen und dort Aufgaben übernehmen, die bisher menschliche Interaktion erforderten.

Robotic Process Automation (RPA)

Robotic Process Automation (RPA, deutsch: Robotergestützte Prozessautomatisierung) ist ein Ansatz zur Geschäftsprozessautomatisierung, bei dem repetitive, manuelle, zeitintensive oder fehleranfällige Tätigkeiten durch sogenannte Softwareroboter (Bots) erlernt und automatisiert ausgeführt werden.³¹

Das Prinzip von Robotic Process Automation ist vergleichbar mit der Automatisierung physischer Arbeit in der Fertigung, wird jedoch auf informationstechnische Arbeitstätigkeiten in indirekten Bereichen von Unternehmen übertragen. Ein besonderes Kennzeichen von RPA ist, dass die Software-Roboter auf bereits bestehende IT-Systeme und IT-Anwendungen aufgesetzt werden und eine vorher menschlich durchgeführte Interaktion mit vorhandenem Benutzer nachahmen.

Abgrenzung zur klassischen Prozessautomatisierung

Im Gegensatz zur klassischen Geschäftsprozessautomatisierung, die meist über Business-Process-Management-Systeme (BPMS) umgesetzt wird, zeichnen sich RPA-Lösungen durch eine einfache, schnelle und kostengünstige Umsetzung aus. Bei der traditionellen Automatisierung müssen neue Schnittstellen eingerichtet und bestehende Anwendungen tiefgreifend angepasst werden, was umfangreiche Tests und einen hohen Entwicklungsaufwand bedeutet. RPA dagegen kann wie ein menschlicher Nutzer direkt über die Benutzeroberfläche arbeiten, ohne zwingend ins Backend eingreifen zu müssen. Dies macht RPA zu einer Art «Low Cost Automation» mit kurzen Amortisationszeiten, die besonders für indirekte Bereiche geeignet ist.

³¹https://de.wikipedia.org/wiki/Robotic_Process_Automation

4 Weitere Automatisierungsformen und Tools

Moderne RPA-Plattformen arbeiten jedoch keineswegs ausschliesslich über die grafische Benutzeroberfläche (GUI). Vielmehr verfolgen sie einen **hybriden Automatisierungsansatz**, der zwei Interaktionsebenen kombiniert: Die UI-basierte Automatisierung simuliert menschliche Aktionen wie Mausklicks, Tastatureingaben und das Auslesen von Bildelementen. Die API-basierte Automatisierung kommuniziert direkt mit dem Backend der Anwendungen über definierte Programmierschnittstellen. Diese Kombination ermöglicht es, für jeden Prozessschritt die optimale Technologie zu wählen: API-Verbindungen bieten höhere Stabilität und Geschwindigkeit, da sie von Änderungen der Benutzeroberfläche unberührt bleiben. GUI-Automatisierung hingegen erschliesst auch Altsysteme ohne moderne Schnittstellen und überbrückt Medienbrüche zwischen nicht integrierten Anwendungen.

Diese technologische Flexibilität macht RPA zu einer wertvollen Brückentechnologie für die digitale Transformation. In heterogenen IT-Landschaften mit gewachsenen Systemstrukturen können Unternehmen Anwendungen über die Oberfläche automatisieren, während moderne Cloud-Dienste über APIs angebunden werden.

Attended und Unattended Automation

In RPA wird grundsätzlich zwischen zwei Arten der Automatisierung unterschieden, die sich in ihrer Interaktion mit menschlichen Nutzern unterscheiden: der beaufsichtigten (Attended) und der unbeaufsichtigten (Unattended) Automatisierung.

Attended Automation

Bei der beaufsichtigten Automatisierung – auch als «Robotic Desktop Automation» bezeichnet – arbeitet der Bot als digitaler Assistent auf dem Arbeitsplatzrechner eines Mitarbeitenden. Der Bot wird durch die Interaktion des Nutzers ausgelöst und übernimmt bestimmte Teilaufgaben innerhalb eines komplexen Arbeitsprozesses. Der Mensch behält die Kontrolle und kann bei Bedarf eingreifen. Diese Form eignet sich besonders für Prozesse, die nicht vollständig automatisierbar sind oder bei denen menschliche Entscheidungen erforderlich bleiben.

Unattended Automation

Bei vollautomatisierten Prozessen arbeitet der Bot vollkommen selbständig und führt nach Aktivierung eines festgelegten Triggers entsprechende transaktionsbasierte Aktivitäten. Ein menschliches Eingreifen ist nicht oder nur in Ausnahmefällen nötig. Diese Vollautomatisierung wird oft in Back-Office-Systemen verwendet und unterstützt das Sammeln, Sortieren und Analysieren grosser Datenmengen innerhalb einer Organisation.

Abbildung 34: Attended versus Unattended Automation

Die Wahl zwischen Attended und Unattended Automation hängt von verschiedenen Faktoren ab: der Komplexität des Prozesses, der Notwendigkeit menschlicher Entscheidungen, dem Transaktionsvolumen und den zeitlichen Anforderungen.

Arbeitsweise und Logiken

RPA-Software arbeitet auf der Benutzeroberflächen-Ebene und interagiert mit bestehenden Systemen wie ein menschlicher Nutzer. Die Bots erkennen Bildelemente, lesen Daten aus ver-

schiedenen Quellen, führen Berechnungen durch und übertragen Informationen zwischen Systemen. Dies geschieht durch eine Kombination verschiedener Technologien: Die Bildschirmerkennung identifiziert UI-Elemente wie Schaltflächen, Eingabefelder und Menüs. Maus- und Tastatursteuerung simuliert menschliche Eingaben. Datenextraktion liest strukturierte und semi-strukturierte Informationen aus Dokumenten, Daten und Bildschirmen.

4 Weitere Automatisierungsformen und Tools

RPA-Plattformen modellieren Prozesse in einzelnen Schritten, die häufig als «Step» oder «Task» bezeichnet werden. Viele Plattformen bedienen sich der Technik von Flussdiagrammen oder anderen Entscheidungsbäumen, um so eine grafische Prozessmodellierung mittels Drag-and-Drop zu ermöglichen. Als solches können sie als No-Code- oder Low-Code-Plattform klassifiziert werden, was bedeutet, dass auch Mitarbeitende ohne Programmierkenntnisse einfache Automatisierungen erstellen können. Prozessschritte greifen auf wiederverwendbare Bausteine aus Bibliotheken zurück, sodass keine Programmierung zur Steuerung von Systemkomponenten notwendig ist.

Moderne RPA-Plattformen kombinieren diese regelbasierte Automatisierung zunehmend mit künstlicher Intelligenz für intelligentere Dokumentenverarbeitung und Entscheidungsfindung. Machine Learning erweitert die Fähigkeiten von Standard-RPA um Texterkennung (OCR), Datenextraktion aus

unstrukturierten Dokumenten und kontextuelle Verarbeitung. Diese Verschmelzung von RPA und KI wird unter dem Begriff «Hyperautomation» oder «Intelligent Process Automation» zusammengefasst und markiert den Übergang von isolierten Automatisierungslösungen hin zu vollständig integrierten, lernfähigen End-to-End-Prozessen.

Der RPA-Markt hat in den vergangenen Jahren ein starkes Wachstum erlebt, wobei sich die anfängliche Euphorie inzwischen etwas normalisiert hat. Branchenbeobachter weisen darauf hin, dass RPA in vielen Fällen als Brückentechnologie dient – eine pragmatische Lösung für Prozesse, die sich aufgrund fehlender Schnittstellen oder historisch gewachsener IT-Landschaften nicht anders automatisieren lassen. Für Unternehmen bleibt die Plattformwahl eine wichtige Entscheidung mit Auswirkungen auf Skalierbarkeit, Wartungsaufwand und Integrationsmöglichkeiten.

Die wichtigsten Plattformen im Überblick

Plattform	Zielgruppe	Besonderes Merkmal
UiPath	Breites Spektrum, starke Community	Umfangreiche Schulungsressourcen (UiPath Academy), visuelle Workflow-Erstellung
Automation Anywhere	Entwicklerorientierte Teams	Cloud-native Architektur, Bot Store mit vorgefertigten Automatisierungen
Blue Prism	Grossunternehmen, regulierte Branchen	Fokus auf Governance und Compliance, On-Premise-Ausrichtung
Microsoft Power Automate	Microsoft-365-Anwender	Tiefe Integration ins Microsoft-Ökosystem, oft in bestehenden Lizenzen enthalten

4 Weitere Automatisierungsformen und Tools

Hinweise zur Plattformwahl

Die Entscheidung für eine RPA-Plattform sollte nüchtern anhand der tatsächlichen Anforderungen getroffen werden. Unternehmen, die bereits stark in Microsoft investiert sind, finden in Power Automate einen niederschweligen Einstieg. Organisationen mit strengen Compliance-Anforderungen – etwa im Finanz- oder Gesundheitswesen – tendieren häufig zu Blue Prism. UiPath und Automation Anywhere bieten breite Funktionalität, wobei die langfristige Kostenentwicklung und Anbieterabhängigkeit kritisch zu prüfen sind.

Grundsätzlich gilt: RPA löst keine Prozessprobleme, sondern automatisiert bestehende Abläufe. Vor der Plattformwahl sollte daher geprüft werden, ob der Prozess nicht durch eine saubere API-Integration oder Prozessoptimierung besser adressiert werden kann.

Der Implementierungsprozess

Die erfolgreiche Implementierung von RPA erfordert einen strukturierten Ansatz, der weit über die reine Technologieeinführung hinausgeht.

Ein bewährtes Vorgehen gliedert die RPA-Implementierung in mehrere Phasen, die aufeinander aufbauen und jeweils spezifische Ziele verfolgen. Am Anfang steht die Identifikation geeigneter Automatisierungskandidaten, gefolgt von einer detaillierten Analyse und einem Proof of Concept (PoC). Erst nach erfolgreicher Validierung folgen die eigentliche Entwicklung, das Testing und schliesslich der produktive Betrieb mit kontinuierlicher Überwachung und Optimierung.



Abbildung 35: Implementierungsprozess (eigene Darstellung)

4 Weitere Automatisierungsformen und Tools

Phase 1: Prozessidentifikation und -auswahl

Für die Grundausswahl potenzieller Prozesse können die «Rule of Five» herangezogen werden, die besagen, dass ein idealer RPA-Kandidat mindestens fünf der folgenden Eigenschaften aufweisen sollte:

1. hohe Transaktionszahlen
2. regelbasierte Entscheidungslogik
3. standardisierte Eingaben und Ausgaben
4. niedrige Ausnahmequote sowie
5. stabile, gut dokumentierte Prozesse.

Phase 2: Proof of Concept

Ein Proof of Concept dient dazu, die grundsätzliche Machbarkeit zu validieren, die gewählte RPA-Plattform in der eigenen Umgebung zu testen und erste Erfahrungen zu sammeln. Der PoC sollte einen repräsentativen, aber überschaubaren Prozess umfassen, der innerhalb weniger Wochen automatisiert werden kann.

Phase 3–4: Entwicklung und Testing

Die eigentliche Entwicklung der RPA-Bots erfordert eine enge Zusammenarbeit zwischen den Entwicklern und den Fachmitarbeitenden, die den zu automatisierenden Prozess kennen. Dabei empfiehlt sich eine agile Vorgehensweise, bei der in kurzen Sprints entwickelt, getestet und verbessert wird.

Das Testing von RPA-Bots ist ein kritischer Erfolgsfaktor: Es sind ausführliche Tests erforderlich, um eine ausreichende Stabilität zu gewährleisten. Dabei müssen nicht nur die Standardfälle, sondern auch alle identifizierten Ausnahmen und Randfälle abgedeckt werden.

- Wie verhält sich der Bot, wenn ein unerwarteter Fehler auftritt?
- Wird der Fehler korrekt protokolliert und eskaliert?
- Kann der Prozess an der richtigen Stelle wieder aufgenommen werden?

Phase 5: Go-Live und Betrieb

Nach erfolgreichem Testing erfolgt die Produktivsetzung des Bots. In der Praxis hat sich ein schrittweises Hochfahren des Transaktionsvolumens bewährt, da dies eine kontrollierte Überwachung der Performance unter realen Bedingungen ermöglicht. Die ersten Tage und Wochen nach dem Go-Live sind typischerweise von intensiver Überwachung geprägt, um auftretende Probleme frühzeitig zu identifizieren. Die für den erfolgreichen Einsatz und Betrieb von RPA notwendige IT-Infrastruktur zählt erfahrungsgemäss zu den am meisten vernachlässigten Aspekten während des Proof of Concepts – sie bildet jedoch eine zentrale Voraussetzung für den produktiven Einsatz.

Monitoring und Wartung

Der Betrieb von RPA-Bots erfordert kontinuierliches Monitoring und regelmässige Wartung. Im Gegensatz zu menschlichen Mitarbeitenden verfügen Bots nicht über die Fähigkeit, selbständig auf Veränderungen in ihrer Umgebung zu reagieren. Bereits geringfügige Änderungen an der Benutzeroberfläche einer Anwendung können dazu führen, dass ein Bot seine Funktion nicht mehr korrekt ausführt. Proaktives Monitoring dient dazu, potenzielle Probleme frühzeitig zu erkennen, bevor sie zu Ausfällen führen.

Die zentrale Steuerungskomponente – in der Fachsprache als Orchestrator bezeichnet – stellt umfangreiche Monitoring-Funktionen bereit, die einen Überblick über alle laufenden und geplanten Bot-Ausführungen ermöglichen. Zu den zentralen Metriken zählen die Erfolgs- und Fehlerquoten der Bot-Ausführungen, die Durchlaufzeiten einzelner Prozesse, die Auslastung der verfügbaren Bot-Kapazitäten sowie die Anzahl der verarbeiteten Transaktionen. Automatische Alerts informieren bei Problemen die zuständigen Personen. Ein strukturiertes Logging ermöglicht die schnelle Diagnose von Fehlern und die Identifikation ihrer Ursachen – mit zunehmender Komplexität und Grösse des Prozesses gewinnt diese Fähigkeit an Bedeutung.

4 Weitere Automatisierungsformen und Tools

Herausforderungen und Erfolgsfaktoren

Das Potenzial von Software-Robotern ist unumstritten. Die Entwicklung und Implementierung von RPA stellt in der Praxis jedoch eine der grössten Herausforderungen dar, die vielfach mit unrealistischen Erwartungen verbunden ist. Häufig gehen Auftraggeber von einer einfachen und schnellen Implementierung sowie exorbitanten Kosteneinsparungen aus. Die Herausforderungen beim Einsatz von RPA lassen sich in vier Hauptkategorien einteilen: technische, fachliche, organisatorische und strategische Herausforderungen

- **Technische Herausforderungen** zeigen sich typischerweise in Form von Störungen nach der Implementierung – verursacht durch zu grosse Datenmengen, Änderungen an zugrundeliegenden Systemen ohne rechtzeitige Kommunikation oder unpassende Regeln in der Entwicklung. Die Kompatibilität mit Bestandssystemen und die notwendige Infrastruktur werden in der Praxis häufig unterschätzt.
- **Fachliche Herausforderungen** entstehen durch mangelndes Wissen über den zu automatisierenden Prozess auf Seiten der Entwickler, durch übersehene Ausnahmen und Sonderfälle sowie durch die Schwierigkeit, implizites Expertenwissen in formale Regeln zu übersetzen.
- **Organisatorische Herausforderungen** manifestieren sich in Widerständen von Mitarbeitenden, die Sorgen vor Arbeitsplatzverlust haben, in mangelnder Abstimmung zwischen IT und Fachabteilung sowie in kurzfristigen Umplanungen bei Bot-Störungen, die Ressourcen binden.
- **Strategische Herausforderungen** umfassen das Phänomen der «Automatisierung von Verschwendung» – also die Automatisierung suboptimaler Prozesse –, die Prozessfragmentierung durch heterogene IT-Systemlandschaften sowie Abhängigkeiten von Anbietern und deren Lizenzmodellen.

Ein besonders häufig zu beobachtender Fehler ist der Versuch, zu komplexe Prozesse zu automatisieren. Viele Verzweigungen und Entscheidungspunkte erhöhen den Entwicklungsaufwand signifikant. Während der Entwicklung können Ausnahmen bei einzelnen Prozessschritten übersehen werden, was nach der Implementierung zu Störungen führt. Klassische RPA ist konzeptionell starr und reagiert empfindlich auf Abweichungen – entsprechend eignen sich vor allem klar strukturierte, regelbasierte Prozesse ohne übermässige Varianten in ihrer Ausführung.

Erfolgsfaktoren für RPA-Projekte

Die Erfolgsfaktoren für RPA-Projekte sind vielfältig und reichen von der sorgfältigen Prozessauswahl über die technische Implementierung bis hin zum organisatorischen Change Management. Eine enge Zusammenarbeit zwischen den Entwicklern und den Fachmitarbeitenden sowie mit dem System ist eine Grundvoraussetzung für erfolgreiche RPA-Projekte.

Die Automatisierung mittels RPA dient nicht dazu, Probleme in den Abläufen zu beheben – wer einen schlechten Prozess automatisiert, erhält einen automatisierten schlechten Prozess. Prozesse erfordern daher vor dem Einsatz von RPA eine gründliche Analyse und gegebenenfalls Optimierung. Der digitale Input muss akkurat und strukturiert sein, da Bots mit unstrukturierten oder qualitativ schlechten Daten nicht effektiv arbeiten können.

Die frühe Einbindung aller Stakeholder in die Kommunikation trägt dazu bei, Ängste abzubauen und Akzeptanz zu schaffen. Automatisierung kann bei Mitarbeitenden Bedenken auslösen – etwa die Sorge vor Arbeitsplatzverlust oder Kontrollverlust über Prozesse. Transparente Kommunikation, Einbindung in Entscheidungsprozesse und gezielte Schulungsmassnahmen erweisen sich als entscheidend für den Abbau von Vorbehalten und die Förderung der Akzeptanz. Das Implementierungsdesign berücksichtigt idealerweise das gesamte betroffene Team, da sich Rollen und organisatorische Strukturen durch die Automatisierung verändern.

4 Weitere Automatisierungsformen und Tools

«Mit einem Proof of Concept, der schnell Erfolge zeigt und damit Vertrauen schafft, sollte begonnen werden. Es sollte eine agile Vorgehensmethodik gewählt werden, in der in Sprints getestet und verfeinert wird.»³²

Eine **tiefgreifende Prozessanalyse** vor der Automatisierung ist unverzichtbar. Dabei stehen mehrere Fragen im Zentrum: Ist der Prozess überhaupt für RPA geeignet? Bietet die Automatisierung einen echten Mehrwert? Was soll mit RPA im Unternehmen erreicht werden? Und welche Ressourcen sind verfügbar – in Bezug auf IT-Struktur, Budget und Know-how? Die ehrliche Beantwortung dieser Fragen vor grösseren Investitionen bildet eine wesentliche Grundlage für den Projekterfolg.

Schliesslich ist eine **realistische Kalkulation von Kosten und ROI** entscheidend. Neben Lizenz- und Implementierungskosten umfasst diese auch Wartung, Skalierung und Prozessmonitoring.

Häufige Fehler vermeiden

Der Beginn mit den komplexesten Prozessen statt mit überschaubaren Quick Wins verzögert frühe Erfolgserlebnisse und gefährdet die Akzeptanz. Die Unterschätzung des Wartungsaufwands für Bots nach der Implementierung führt zu Ressourcenengpässen im laufenden Betrieb. Zudem erweist sich RPA nicht für jeden Anwendungsfall als optimale Lösung – manche Prozesse sind für andere Automatisierungsansätze oder eine grundlegende Neugestaltung besser geeignet.

4.3 Cloud- und Toolbasierte Automatisierung

Grundlagen der Cloud-Automatisierung

Cloudbasierte No-Code-Automatisierungstools bilden eine eigenständige Kategorie innerhalb der Prozessautomatisierung. Sie ermöglichen die Automatisierung von Geschäftsprozessen ohne Programmierkenntnisse und erfordern im Vergleich zu RPA-Implementierungen deutlich geringere Investitionen. Insbesondere für kleine und mittlere Unternehmen stellen No-Code-Plattformen häufig einen niederschweligen Einstiegspunkt in die Automatisierung dar.

Das Funktionsprinzip dieser Plattformen basiert auf der Verbindung verschiedener Cloud-Anwendungen und der Automatisierung des Datenaustauschs zwischen ihnen. Prozesse, die zuvor eine manuelle Datenübertragung von einem System in ein anderes erforderten – etwa die Erfassung von Kundenanfragen aus einem Webformular im CRM-System – werden durch die Automatisierungsplattform übernommen. Das Grundprinzip folgt dabei einem einheitlichen Muster: Ein auslösendes Ereignis – der sogenannte Trigger – initiiert eine oder mehrere Aktionen in verbundenen Anwendungen.

Trigger + Aktion = Automation

Jede Automatisierung besteht aus einem auslösenden Ereignis (Trigger) und einer oder mehreren daraus resultierenden Aktionen. Beispiel: «Wenn eine neue E-Mail mit Anhang eingeht (Trigger), dann speichere den Anhang in einem Cloud-Ordner und erstelle eine Aufgabe im Projektmanagement-Tool (Aktionen).»

Der entscheidende Unterschied zu RPA liegt im technischen Ansatz: Während RPA-Bots die Benutzeroberfläche von Anwendungen simulieren – also gewissermassen «von aussen» auf Programme zugreifen –, nutzen Cloud-Automatisierungstools die programmierten Schnittstellen (APIs) der

³²Quelle: AT Kearney & Arvato (2018): RPA Erfolgsfaktoren

4 Weitere Automatisierungsformen und Tools

Anwendungen. Dies macht sie in der Regel stabiler, da sie nicht von Änderungen an der Benutzeroberfläche betroffen sind, aber es setzt voraus, dass die zu verbindenden Anwendungen entsprechende Schnittstellen anbieten. Die meisten modernen Cloud-Anwendungen – von Buchhaltungssoftware über CRM-Systeme bis zu Office-Produkten – bieten heute solche Schnittstellen an.

Abgrenzung zu RPA

RPA ist besonders dann sinnvoll, wenn mit veralteten Systemen gearbeitet werden muss, die keine modernen Schnittstellen bieten, oder wenn komplexe Desktop-Anwendungen automatisiert werden sollen. Cloudbasierte No-Code-Tools sind hingegen die wirtschaftlichere Wahl, wenn die verwendeten Anwendungen bereits cloudbasiert sind und APIs anbieten. In der Praxis bedeutet dies: Wer bereits mit Microsoft 365, Google Workspace, Salesforce, HubSpot oder ähnlichen modernen Plattformen arbeitet, kann mit No-Code-Tools oft schneller und günstiger Ergebnisse erzielen als mit RPA.

4.4 Praktische Anwendungsfälle für KMU

Die Anwendungsfelder cloudbasierter Automatisierung in KMU erstrecken sich über nahezu alle Unternehmensbereiche. Die technische Umsetzung erfolgt auf unterschiedlichen Wegen. Mit Power Automate lassen sich beispielsweise Genehmigungsworkflows direkt in Microsoft Teams integrieren: Der Antragsteller erfasst die Anfrage über ein Formular, der Vorgesetzte erhält automatisch eine Benachrichtigung mit Genehmigungsoptionen, und nach erfolgter Genehmigung wird der Eintrag automatisch in die relevanten Systeme übertragen.

Ein häufiger Anwendungsbereich ist die Automatisierung personaladministrativer Prozesse. Urlaubsanträge lassen sich digital erfassen und über automatische Genehmigungsworkflows durch Vorgesetzte freigeben. Die mobile Zeiterfassung ermöglicht eine direkte Integration in Lohnabrechnungssysteme ohne manuelle Datenübertragung. Bei Krankmeldungen kann der digitale Prozess den

automatischen Abruf der elektronischen Arbeitsunfähigkeitsbescheinigung (eAU) bei der Krankenkasse einschliessen. Spesenabrechnungen erfolgen über den Upload von Belegen per App mit automatischer Weiterleitung an die Buchhaltung. Stammdaten wie Adresse, Bankverbindung oder Kontaktdaten können von Mitarbeitenden selbstständig aktualisiert werden, während Gehaltsabrechnungen digital bereitgestellt und archiviert werden.

Die Integration von Automatisierungstools mit bestehenden ERP-Systemen stellt für KMU ein zentrales Thema dar. Moderne ERP-Systeme wie SAP Business One, Microsoft Dynamics oder Bexio verfügen in der Regel über API-Schnittstellen, die eine Anbindung an Automatisierungsplattformen ermöglichen. Dadurch lassen sich beispielsweise Daten aus Webformularen automatisch in das ERP-System übernehmen, Auftragsbestätigungen bei Bestelleingang versenden oder Lagerbestände zwischen verschiedenen Systemen synchronisieren.

Urlaubsanträge Digitale Beantragung mit automatischem Genehmigungsworkflow durch Vorgesetzte	Zeiterfassung Mobile Erfassung von Arbeitszeiten mit direkter Integration in Lohnabrechnung	Krankmeldungen Digitaler Prozess mit automatischem Abruf der eAU bei der Krankenkasse
Spesenabrechnungen Upload von Belegen per App, automatische Weiterleitung an Buchhaltung	Stammdaten Selbstständige Aktualisierung von Adresse, Bankverbindung, Kontaktdaten	Gehaltsabrechnungen Digitale Bereitstellung und Archiv aller Lohndokumente

Abbildung 36: Typische Anwendungsfelder von Automatisierungstools³³

Die erfolgreiche Einführung von Automatisierungslösungen erfordert einen strukturierten Ansatz, der technische und organisatorische Aspekte gleichermaßen berücksichtigt.

³³Quelle: ATKearney & Arvato (2018): RPA Erfolgsfaktoren

4 Weitere Automatisierungsformen und Tools

Vorgehensmodell zur Einführung cloud-basierter Automatisierung

Phase 1: IST-Analyse und Prozessauswahl

Am Anfang steht die systematische Bestandsaufnahme der bestehenden Prozesse. Dabei sind folgende Fragestellungen relevant: Welche Tätigkeiten fallen regelmässig an? Welche erweisen sich als zeitaufwändig, fehleranfällig oder besonders repetitiv? Die Priorisierung erfolgt nach dem Automatisierungspotenzial, wobei hohe Wiederholungsfrequenz, klare Regelbasierung sowie standardisierte Ein- und Ausgaben als positive Indikatoren gelten.

Phase 2: Identifikation eines Quick Wins

Ideale Automatisierungen zeichnen sich durch eine kurze Implementierungsdauer (Tage statt Wochen), einen unmittelbar spürbaren Nutzen sowie Relevanz für einen grösseren Mitarbeiterkreis aus. Typische Beispiele umfassen automatische E-Mail-Benachrichtigungen, Formular-zu-Tabelle-Workflows oder einfache Freigabeprozesse.

Phase 3: Proof of Value (PoV)

Im Unterschied zum klassischen «Proof of Concept», der primär die technische Machbarkeit demonstriert, zielt der Proof of Value auf den Nachweis des konkreten Nutzens ab:

- Wie viel Zeit wird eingespart?
- Wie hat sich die Fehlerquote verändert?

Diese Kennzahlen bilden die Grundlage für die Akzeptanz im Unternehmen und die Freigabe weiterer Investitionen.

Phase 4: Schrittweise Ausrollung

Nach erfolgreichem PoV folgt die Automatisierung weiterer Prozesse in iterativer Vorgehensweise: Identifikation des nächsten Kandidaten, Umsetzung, Messung, Erkenntnisgewinn. Mit jeder Iteration wächst das Know-how im Team und die Automatisierungslandschaft wird umfangreicher – bei gleichzeitig kontrolliertem und nachhaltigem Wachstum.

Phase 5: Schulung und Dokumentation

Die Befähigung der Mitarbeitenden stellt einen wesentlichen Erfolgsfaktor dar. Fundiertes Verständnis der eingesetzten Tools ermöglicht nicht nur die Wartung bestehender Automatisierungen, sondern auch die Entwicklung neuer Anwendungsfälle. Eine sorgfältige Dokumentation aller Automatisierungen – Funktionsumfang, beteiligte Systeme, Verantwortlichkeiten – erleichtert Wartung und Weiterentwicklung erheblich.

4 Weitere Automatisierungsformen und Tools

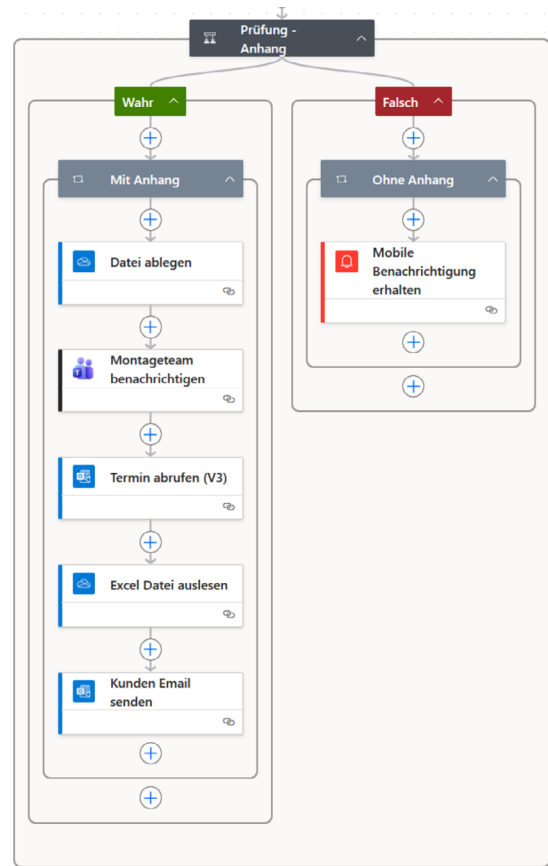
Case Study: Automatisierung von Mail-Anlagen und Kundenprozesse

Sie erhalten von einem Produzenten täglich, jeweils am Abend, eine E-Mail mit einer Excel-Datei, welche in einer einfachen Tabelle auflistet, welche Produkte produziert wurden.

Diese Liste wurde bis anhin manuell in einem Ordner gespeichert und dann in eine Gesamtliste integriert, damit die Sales-Mitarbeitenden die Kunden informieren und einen Liefer- und Montagetermin festlegen konnten, nachdem sie mit dem Montageteam die Auslastung geprüft hatten.

Wir prüfen, ob eine E-Mail von XY eintrifft und ob diese einen Anhang hat. Falls WAHR, wird diese in einem Ordner abgelegt, falls FALSCH, wird der Produzent informiert, dass die Anlage fehlt. Zusätzlich wird die Anzahl der WAHR / FALSCH Meldungen in einer Logdatei gespeichert; so können wir auch die Qualität des Versands prüfen.

Ausserdem lesen wir die Kundennummer aus der Liste, gruppieren die Bestellung und senden eine E-Mail an den Kunden; dabei wird auch das Montageteam benachrichtigt, damit diese die Planung vornehmen kann.



Klein anfangen

Lieber einen Prozess richtig automatisieren als zehn nur halb. Starten Sie mit überschaubaren Quick Wins, die schnell Erfolge zeigen und ein Momentum schaffen.

Nutzen messen

Definieren Sie vor der Automatisierung, wie Sie den Erfolg messen werden. Zeitersparnis, Fehlerreduktion, Durchlaufzeit – nur was gemessen wird, kann verbessert werden.

Menschen mitnehmen

Transparente Kommunikation über Ziele und Vorteile. Mitarbeitende sollten verstehen, dass Automatisierung sie entlastet, nicht ersetzt.

Prozesse erst optimieren

Automatisieren Sie keine schlechten Prozesse. Analysieren und optimieren Sie zuerst, dann automatisieren Sie den verbesserten Ablauf.

Know-how aufbauen

Investieren Sie in Schulungen. Interne Kompetenz reduziert Abhängigkeiten von externen Dienstleistern und beschleunigt zukünftige Projekte.

Kontinuierlich verbessern

Automatisierung ist kein einmaliges Projekt, sondern ein fortlaufender Prozess. Regelmässiges Review und Optimierung sind Teil des Betriebs.

Abbildung 37: Erfolgsfaktoren eines Automatisierungsprojekts (eigene Darstellung)

4 Weitere Automatisierungsformen und Tools

Ein besonders kritischer Erfolgsfaktor ist das **Change Management**. Automatisierung verändert Arbeitsabläufe und kann bei Mitarbeitenden Bedenken auslösen – etwa die Sorge vor Arbeitsplatzverlust oder Kontrollverlust über die eigenen Prozesse. Gerade in Zeiten der digitalen Transformation muss Change Management transparent sein und mit einer klaren Vision den Weg aufzeigen. Führungskräfte fungieren dabei als Bindeglied zwischen Strategie und Belegschaft. Eine offene Dialogkultur, in der Fragen, Ideen und auch Bedenken gehört und diskutiert werden, ist wichtig für den Erfolg. Mitarbeitende sollten nicht nur über die Transformation informiert, sondern aktiv in die Entwicklung eingebunden werden.

Risiken und Abhängigkeiten

Neben den Chancen bringt die Cloud-Automatisierung auch Risiken und Abhängigkeiten mit sich, die bei der Umsetzung berücksichtigt werden sollten.

Anbieterabhängigkeit

Mit jedem Automatisierungstool entsteht eine Abhängigkeit vom Anbieter. Preisänderungen, Funktionseinschränkungen oder im schlimmsten Fall die Einstellung des Dienstes können erhebliche Auswirkungen haben. Eine sorgfältige Dokumentation aller Workflows erleichtert im Notfall die Migration.

Datenschutz und Sicherheit

Cloud-Dienste verarbeiten oft sensible Unternehmensdaten. Es muss sichergestellt sein, dass die gewählten Plattformen DSGVO- bzw. DSGVO-konform arbeiten und angemessene Sicherheitsmassnahmen bieten. Serverstandorte, Verschlüsselung und Zugriffskontrollen sollten geprüft werden.

Wartungsaufwand

Automatisierungen müssen gepflegt werden. Wenn sich verbundene Anwendungen ändern – etwa durch Updates oder neue Versionen – können Workflows brechen. Regelmässiges Monitoring und Ressourcen für Wartung sind einzuplanen.

Know-how-Risiko

Wenn das Wissen über Automatisierungen nur bei einer Person liegt, entsteht ein Risiko. Dokumentation, Schulung mehrerer Mitarbeitender und die Vermeidung von «Wissensinseln» sind wichtige Gegenmassnahmen.

Ein oft unterschätzter Aspekt ist die sogenannte «Automatisierung von Verschwendung». Wer einen ineffizienten oder fehlerhaften Prozess automatisiert, erhält einen automatisierten ineffizienten Prozess. Deshalb sind die kritische Analyse und gegebenenfalls die Optimierung eines Prozesses vor der Automatisierung wichtig. Die Frage «Sollten wir diesen Prozess überhaupt so durchführen?» muss vor der Frage «Wie können wir diesen Prozess automatisieren?» beantwortet werden.

4.5 Herausforderungen und Erfolgsfaktoren

Die Einführung von Datenautomatisierung stellt Schweizer KMU vor vielfältige Herausforderungen, die weit über rein technische Fragestellungen hinausgehen. Während die vorangegangenen Abschnitte die verschiedenen Automatisierungstechnologien von ETL-Prozessen über RPA bis hin zu cloudbasierten Lösungen im Detail beleuchtet haben, widmet sich dieser abschliessende Abschnitt den übergreifenden Erfolgsfaktoren und den typischen Stolpersteinen, die über Gelingen oder Scheitern von Automatisierungsinitiativen entscheiden.

Studien zeigen konsistent, dass ein erheblicher Teil aller Digitalisierungs- und Automatisierungsprojekte die ursprünglich gesteckten Ziele nicht erreicht. Studien zeigen wiederholt, dass der Grossteil gescheiterter Automatisierungsprojekte nicht an der Technologie scheitert, sondern an mangelnder Vorbereitung – insbesondere fehlender Prozesstransparenz und unzureichender Datenqualität. Diese Erkenntnis verdeutlicht, dass der Erfolg von Automatisierungsvorhaben massgeblich von einer systematischen Herangehensweise, der Einbindung der Mitarbeitenden und der Berücksichtigung regulatorischer Rahmenbedingungen abhängt. Für

4 Weitere Automatisierungsformen und Tools

Schweizer Unternehmen kommen dabei spezifische Compliance-Anforderungen hinzu, insbesondere das seit September 2023 geltende neue Datenschutzgesetz (DSG) sowie branchenspezifische Regulierungen wie die FINMA-Vorgaben für den Finanzsektor.

Die Erfahrung aus zahlreichen Automatisierungsprojekten zeigt ein wiederkehrendes Muster: Technische Hürden lassen sich in der Regel überwinden. Projekte scheitern häufiger an unklaren Prozessen, mangelnder Datenqualität oder fehlender Einbindung der Betroffenen. Eine gründliche Vorbereitung – insbesondere die Dokumentation und Analyse der bestehenden Abläufe – erhöht die Erfolgswahrscheinlichkeit erheblich.

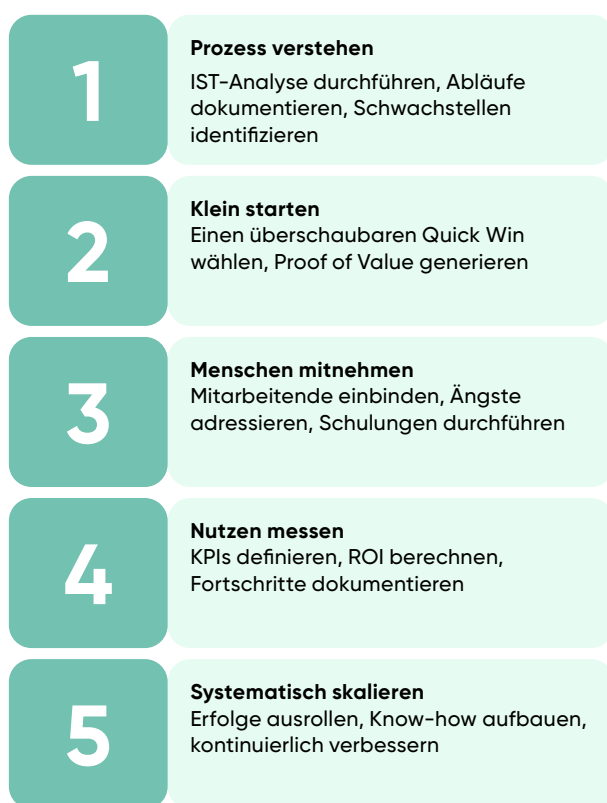


Abbildung 38: Die fünf kritischen Erfolgsfaktoren (eigene Darstellung)

1. Prozess verstehen vor Automatisierung

Der häufigste Fehler bei Automatisierungsprojekten ist der Versuch, schlecht verstandene oder ineffiziente Prozesse zu automatisieren. Dadurch wird lediglich die Verschwendung beschleunigt, nicht die Wertschöpfung. Eine gründliche IST-Analyse umfasst die Dokumentation aller Prozessschritte, die Identifikation von Medienbrüchen und manuellen Eingriffen, das Erfassen von Ausnahmen und Sonderfällen sowie das Verstehen der beteiligten Systeme und Datenflüsse. Erst wenn der Prozess vollständig transparent ist, können sinnvolle Automatisierungspotenziale identifiziert werden.

2. Klein starten mit Quick Wins

Ein funktionierender automatisierter Prozess ist mehr wert als zehn halbfertige. Der Einstieg sollte mit einem überschaubaren Prozess erfolgen, der regelmässig vorkommt, klar definierte Regeln hat und aktuell spürbar Zeit kostet. Kritisch ist dabei, keinen geschäftskritischen Kernprozess als Piloten zu wählen, sondern einen Bereich, in dem ein Fehler keine schwerwiegenden Folgen hätte. Ein typischer Kandidat ist etwa die automatische Kategorisierung und Weiterleitung von Kontaktformular-Anfragen oder die strukturierte Ablage von Eingangsrechnungen.

3. Menschen aktiv einbinden

Automatisierungsprojekte sind immer auch Change-Management-Projekte. Wenn Mitarbeitende nicht frühzeitig informiert und eingebunden werden, entstehen Unsicherheiten und Widerstände, die Projekte ausbremsen oder zum Scheitern bringen. Studien zeigen, dass Unternehmen mit strukturierten Schulungsprogrammen eine dreifach höhere Akzeptanzrate erreichen. Die Betroffenen müssen zu Beteiligten gemacht werden: Sie kennen die Prozessdetails und Ausnahmefälle, an denen Standardlösungen scheitern. Eine transparente Kommunikation darüber, dass Automatisierung nur ein Werkzeug ist, das die Arbeit erleichtern und nicht ersetzen soll, ist dabei essenziell.

4. Nutzen messbar machen

Ohne definierte Erfolgskriterien lässt sich weder der Fortschritt nachvollziehen, noch der Return on

4 Weitere Automatisierungsformen und Tools

Investment belegen. Relevante Kennzahlen umfassen die eingesparte Bearbeitungszeit pro Vorgang, die Reduktion der Fehlerquote, die Verkürzung von Durchlaufzeiten sowie die Kosteneinsparungen durch wegfallende manuelle Tätigkeiten. Diese Metriken sollten vor Projektbeginn erhoben werden, um einen belastbaren Vergleichswert zu haben. Erst mit den richtigen Metriken wird sichtbar, was funktioniert und wo nachjustiert werden muss.

5. Systematisch skalieren

Nach dem erfolgreichen Pilotprojekt gilt es, die gewonnenen Erkenntnisse systematisch auf weitere Prozesse zu übertragen. Dies erfordert den Aufbau von internem Know-how, die Dokumentation von Best Practices und die schrittweise Erweiterung des Automatisierungsportfolios. Wichtig ist dabei, nicht zu schnell zu viel auf einmal anzugehen. Die Erfahrung zeigt, dass viele Unternehmen, die diesen schrittweisen Ansatz wählen, innerhalb von zwölf Monaten mehrere stabile Automatisierungen entwickeln, während andere noch an ihrem ersten großen Projekt arbeiten.



Abbas Tutcouglu

5.1 Was ist künstliche Intelligenz?

Wenn heute ein Begriff in eine Suchmaschine eingegeben wird, schlägt das System häufig automatisch passende Ergänzungen vor. Werden Fotos in einen cloudbasierten Fotospeicher hochgeladen, erkennt das System typische Bildinhalte und versieht sie mit Stichworten. Bei der Beantragung einer Kreditkarte wird auf Basis eingegebener Informationen eine Einschätzung der Kreditwürdigkeit vorgenommen. Und bei digitalen Abonnements erhalten Nutzerinnen und Nutzer regelmässig Empfehlungen, die auf ihr individuelles Verhalten zugeschnitten sind.

Diese Beispiele stammen aus unterschiedlichen Anwendungsfeldern, beruhen jedoch auf demselben Grundprinzip: **Computermodelle analysieren historische Daten, um darin Muster zu erkennen und diese auf neue Situationen anzuwenden.**

Konkret bedeutet dies unter anderem:

- Suchmaschinen lernen aus vergangenen Suchanfragen, welche Begriffe häufig gemeinsam auftreten
- Bilderkennungsmodelle werden anhand klassifizierter Bilder trainiert und erkennen visuelle Merkmale
- Kreditwürdigkeitsmodelle analysieren historische Finanz- und Verhaltensdaten
- Anbieter von Abonnements segmentieren Nutzergruppen auf Basis gemeinsamer Eigenschaften

Im Kern imitiert künstliche Intelligenz damit menschliche Entscheidungs- oder Zuordnungsprozesse, indem sie aus vergangenen Beobachtungen Regeln, Zusammenhänge oder Wahrscheinlichkeiten ableitet. Diese Regeln werden nicht explizit programmiert, sondern implizit aus Daten gelernt.

Künstliche Intelligenz (KI, engl. Artificial Intelligence, AI) beschreibt den Vorgang, **aus Daten Muster zu erkennen und diese erlernten Muster auf neue Daten anzuwenden**, um Einschätzungen oder Vorhersagen zu erzeugen. Dieser Vorgang wird als Lernen oder Training bezeichnet. Da dieses Training rechnergestützt erfolgt, spricht man von maschinellem Lernen (Machine Learning).

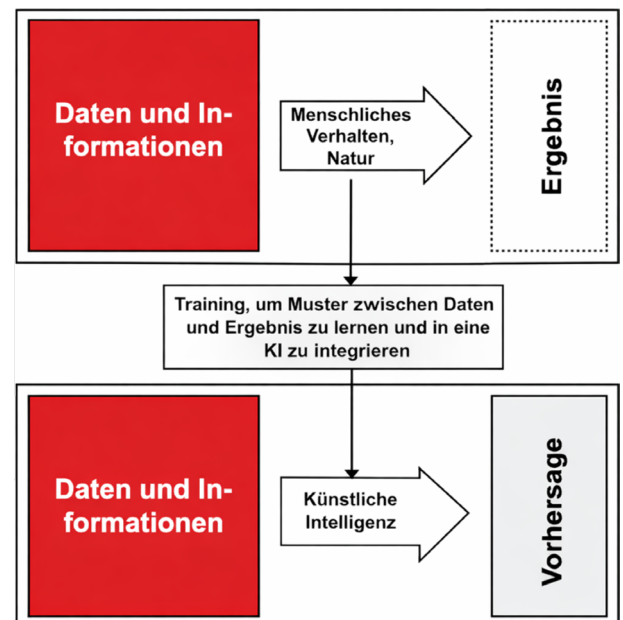


Abbildung 39: KI als Musterspiegel menschlicher oder anderweitiger Entscheidungsfindung (eigene Darstellung)

Im alltäglichen Sprachgebrauch werden zudem Begriffe wie Advanced Analytics oder Predictive Analytics verwendet. Diese sind jedoch nicht deckungsgleich mit KI, sondern bezeichnen spezifische Anwendungsformen oder Teilbereiche.

In den letzten Jahren hat sich dieses Grundprinzip deutlich weiterentwickelt. Neben klassischen, auf einzelne Aufgaben spezialisierten KI-Modellen sind **generative KI-Systeme** (auch GenAI genannt) entstanden. Diese Systeme beschränken sich nicht darauf, Daten zu analysieren oder ein einzelnes Ergebnis vorherzusagen, sondern sind in der Lage,

5 Künstliche Intelligenz

neue Inhalte zu erzeugen, etwa Texte, Tabellen, Bilder oder Zusammenfassungen.

Eine zentrale Rolle spielen dabei Sprachmodelle, auch genannt **Large Language Models (LLMs)**. Diese Modelle wurden auf sehr grossen Textmengen trainiert und sind darauf spezialisiert, sprachliche Muster besonders gut zu erfassen. Auch hier basiert die Funktionsweise nicht auf einem menschlichen Sprachverständnis, sondern auf statistischer Mustererkennung: LLMs lernen, welche Wörter, Satzteile oder Bedeutungen typischerweise gemeinsam auftreten, und nutzen diese Muster, um Texte kontextbezogen fortzuführen oder neu zu erzeugen. Moderne Systeme kombinieren diese sprachbasierten Fähigkeiten zunehmend mit anderen Datenarten, etwa Bildern oder strukturierten Informationen.

Für das Finanz- und Rechnungswesen ergibt sich daraus ein spürbarer Wandel. KI-Systeme sind nicht mehr ausschliesslich im Hintergrund tätig, sondern werden zunehmend zu **interaktiven Werkzeugen**, die Fachpersonen bei Analysen, Einschätzungen, Dokumentation und Entscheidungsprozessen direkt unterstützen. Ein grundlegendes Verständnis der

Funktionsweise und Grenzen künstlicher Intelligenz ist daher eine zentrale Voraussetzung für ihren verantwortungsvollen Einsatz.

Wichtige Begriffe aus diesem Abschnitt: Künstliche Intelligenz (KI), Artificial Intelligence (AI), Mustererkennung, Lernen, Training, Vorhersage, maschinelles Lernen (Machine Learning), generative KI, Large Language Models (LLMs), Data Mining, Predictive Analytics.

Mustererkennung erklärt an einem einfachen Beispiel

Betrachten wir ein einfaches Beispiel: Angenommen, wir vermarkten einen Kredit und möchten diesen nur an Personen vergeben, bei denen davon ausgegangen werden kann, dass sie den Kredit zurückzahlen. Für jeden Bewerber stehen uns drei Informationen zur Verfügung:

- ob die Person bereits Kunde bei uns ist
- ob die Person die Farbe Orange mag
- wie oft die Person pro Jahr spazieren geht

Zusätzlich wissen wir aus einem historischen Datensatz, wann der Kredit zurückgezahlt wurde:

Bereits Kunde?	Mag Orange?	Spaziergänge pro Jahr	Kredit zurückgezahlt?
Ja	Nein	3	Ja
Nein	Nein	9	Nein
Nein	Nein	10	Ja
Ja	Ja	5	Ja
Nein	Ja	9	Nein
Ja	Nein	15	Ja
Ja	Ja	10	Ja
Nein	Ja	11	Nein
Nein	Nein	8	Nein
Nein	Nein	11	Ja

5 Künstliche Intelligenz

Versuchen wir nun, Zusammenhänge zwischen den drei Informationen und dem Ergebnis Kredit zurückgezahlt zu erkennen, lassen sich relativ schnell einfache Muster identifizieren.

- Ist der Bewerber bereits Kunde, wurde der Kredit in allen Fällen zurückgezahlt.
- Ist der Bewerber kein Kunde und mag Orange, wurde der Kredit nie zurückgezahlt.
- Ist der Bewerber kein Kunde, mag Orange nicht und geht mind. 10x pro Jahr spazieren, wurde der Kredit zurückgezahlt.

Dieses vereinfachte und zugegeben realitätsferne Beispiel verdeutlicht, was Mustererkennung bedeutet. Genau diesen Prozess automatisiert KI: Sie identifiziert Zusammenhänge in historischen Daten und nutzt diese, um auf neue, bisher unbekannte Fälle zu schließen.

Möchten wir beispielsweise einschätzen, ob ein neuer Bewerber, der kein Kunde ist, Orange nicht mag und zwölf Spaziergänge pro Jahr macht, den Kredit zurückzahlen wird, kann das trainierte Modell diese Einschätzung auf Basis der zuvor erkannten Muster vornehmen und zu dem Schluss gelangen, dass der Kredit mit hoher Wahrscheinlichkeit zurückgezahlt wird.

In der Praxis sind die zugrundeliegenden Daten, Zusammenhänge und Entscheidungsregeln deutlich komplexer. Gerade in solchen Situationen kann KI den Menschen unterstützen, indem sie Muster erkennt, die manuell nur mit hohem Aufwand oder gar nicht identifizierbar wären.

Dieses Grundverständnis bildet die Basis für die nachfolgenden Abschnitte, in denen die verschiedenen Formen der KI sowie deren Einsatzmöglichkeiten und Grenzen im Finanz- und Rechnungswesen systematisch betrachtet werden. Möglicherweise werden die Teilnehmenden dieses Kurses nicht ihre eigenen KI-Modelle bauen. Dennoch ist das Wissen, wie diese Modelle gebaut sind, notwendig, um unbegründete Ängste vor KI aus dem Weg zu räumen und sie bestmöglich nutzen zu können.

Wichtige Begriffe aus diesem Abschnitt: Trainingsdaten, historische Daten, Merkmale (Features), Ergebnis, Modell, Mustererkennung, Vorhersage

5.2 Die verschiedenen Formen der KI

Die bisherigen Abschnitte zeigen das Grundprinzip: KI erkennt Muster in Daten und wendet diese «Musterkenntnis» auf neue Fälle an. In der Praxis unterscheiden sich KI-Anwendungen jedoch darin, welche Daten vorliegen (Text, Zahlen, Bilder etc.), ob ein Ergebnis vorgegeben ist oder nicht und welche Art von Ergebnis erzeugt werden soll (Ja/Nein, eine Zahl, Text etc.).

Überwachtes und unüberwachtes Lernen

Eine grundlegende Unterscheidung betrifft die Trainingsdaten. Beim **überwachten Lernen (supervised learning)** bestehen die Trainingsdaten aus Merkmalen (**Features**) als Eingabe, und einem bekannten Ergebnis (**Labels**).

Das Training ist dabei «überwacht» im Sinne von: Für historische Fälle ist das gewünschte Ergebnis bereits bekannt. Das Modell lernt, den Zusammenhang zwischen Features und Label nachzubilden, um für neue Fälle eine Vorhersage zu ermöglichen.

Beispiel: Für die Prognose des Zahlungsverhaltens von Debitoren umfassen Features z.B. Zahlungsfristen, Mahnstufen, Zahlungsverlässlichkeit, Umsatzhistorie oder Brancheninformationen, während das Label z.B. Zahlungsverzug: Ja oder Nein (Klassifikation) oder die erwartete Verzugsdauer in Tagen (Regression) sein könnte.

Beim **unüberwachten Lernen (unsupervised learning)** gibt es kein vorgegebenes Label. Die Trainingsdaten enthalten nur Features, und das Modell sucht selbst nach Strukturen, Mustern oder Auffälligkeiten. Typische Formen sind:

5 Künstliche Intelligenz

- **Clustering / Segmentierung:** z.B. Kostenstellen oder Kunden in Gruppen mit ähnlichen Eigenschaften einteilen, um Muster in Kostenstrukturen oder Verhaltensprofilen sichtbar zu machen
- **Anomaly Detection (Ausreissererkennung):** z.B. ungewöhnliche Buchungen, atypische Kombinationen aus Konto, Kostenstelle und Betrag oder auffällige Spesenabrechnungen identifizieren
- **Association Rules (Zusammenhangsregeln):** z.B. häufig gemeinsam auftretende Buchungsmuster oder Kombinationen erkennen, die als Plausibilitätscheck oder zur Standardisierung von Kontierungen genutzt werden können.

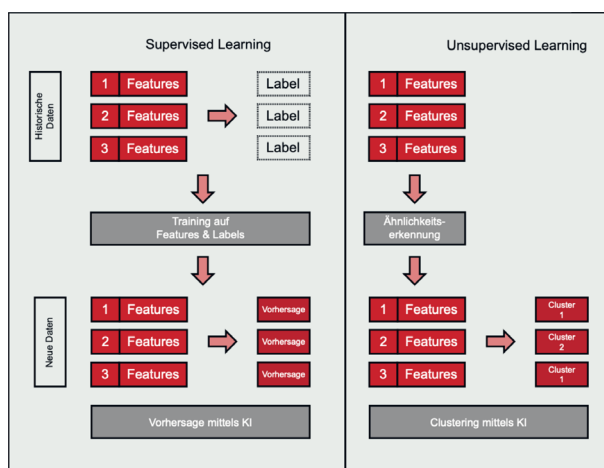


Abbildung 40: Unterschied zwischen supervised und unsupervised learning (eigene Darstellung)

Klassifikation und Regression

Innerhalb des überwachten Lernens wird häufig nach der Art des Labels unterschieden:

Von **Klassifikation (classification)** spricht man, wenn das Label aus Klassen besteht, z.B.:

- *Auffällig / unauffällig*
- *Kreditwürdig: Ja / Nein*
- *Beleg korrekt zugeordnet: Ja / Nein*
- oder auch mehrere Klassen, etwa *Kostenart A / B / C* (Multiclass).

Von **Regression** spricht man, wenn das Label ein kontinuierlicher Zahlenwert ist, z.B.:

- Prognose von Cashflow oder Umsatz
- Schätzung der Höhe einer Rückstellung
- Erwarteter Forderungsausfall in CHF

Analytische KI und generative KI

Eine weitere praxisnahe Einteilung unterscheidet zwischen Modellen, die primär analysieren und vorhersagen, und Modellen, die Inhalte erzeugen.

Analytische KI umfasst insbesondere Systeme, die Klassen, Wahrscheinlichkeiten oder Zahlenwerte liefern. Sie ist häufig dort stark, wo strukturierte Daten vorliegen, und eine klare Zielgrösse definiert ist, z.B. bei Prognosen, Risikoindikatoren oder Auffälligkeitsanalysen.

Generative KI (GenAI) erzeugt hingegen neue Inhalte auf Basis gelernter Muster, typischerweise in Sprache, aber auch in weiteren Darstellungsformen. Im Controlling zeigt sich der Nutzen u. a. dort, wo Ergebnisse erklärt, zusammengeführt oder formuliert werden müssen, z.B.:

- Antwort auf eine Frage zu einer Rechnung
- Entwurf von Management-Kommentaren zu Abweichungen
- Zusammenfassung von Reportings oder Sitzungsnotizen
- Formulierungsvorschläge für Abschlussanhänge oder interne Richtlinien

5 Künstliche Intelligenz

- Unterstützung bei der Strukturierung von Analysen (z. B. Fragenkataloge)
- Extraktion von Schlüsselinformationen aus längeren Texten (z. B. der Kundennummer aus einem Briefscan). Obwohl dieses Beispiel auf den ersten Blick nicht als «Generierung» erscheint, fällt es darunter, da mithilfe von Kontexterkennung relevante Attribute aus unstrukturiertem Text abgeleitet und befüllt werden.

Wichtig für die Einordnung: Oft ist eine Kombination generativer und analytischer KI sinnvoll. So kann eine generative KI aus langen Texten, wie etwa Bankbelegen, wichtige Daten entnehmen, um somit die Features für ein analytisches Modell, z. B. zur automatischen Zuordnung von Transaktionen zu Forderungen, zu befüllen.

KI nach Art der Eingabedaten

Eine weitere Einteilung orientiert sich an der Art der Daten, die eingespeist werden:

- **Tabellarisches maschinelles Lernen (tabular machine learning)** nutzt strukturierte Tabellen, wie sie im Finanzwesen typisch sind.
- **Texterkennung und Sprachverarbeitung** arbeitet mit Texten, etwa Buchungstexten, Verträgen, E-Mails oder Kommentaren. Sie kann z. B. Texte klassifizieren, extrahieren, zusammenfassen oder inhaltlich strukturieren.
- **Bilderkennung** nutzt Bilddaten, z. B. bei der automatisierten Verarbeitung von Belegen oder Dokumenten. Häufig steht hier die Extraktion relevanter Informationen (z. B. Rechnungsnummer, Betrag, Lieferant) im Vordergrund.

Generative KI kombiniert in diesem Kontext häufig mehrere Datentypen. So können sog. multimodale Systeme, sowohl Text (zum Beispiel eine Frage zu einer Rechnung), wie auch Dokumente (z. B. die Rechnung selbst) aufnehmen und daraufhin eine Antwort generieren.

Wichtige Begriffe aus diesem Abschnitt: Überwachtes Lernen (Supervised Learning), unüberwachtes Lernen (Unsupervised Learning), Features (Merkmale), Labels, Klassifikation, Regression, Clustering/

Segmentierung, Anomaly Detection (Ausreissererkennung), analytische KI, generative KI (GenAI).

5.3 Entstehung von KI-Modellen

In den bisherigen Abschnitten haben wir verschiedene Beispiele kennengelernt, bei denen KI auf Basis historischer Daten Muster erkennt und diese Muster auf neue Daten anwendet, um Vorhersagen oder Einschätzungen zu treffen. Doch wie entsteht ein solches Modell? Wie funktioniert dieses Training, bei dem ein Modell Muster erkennt und diese auf neue Beispiele anwendet? Und wie kommt es dazu, dass ein KI-Modell – etwa eine Klassifikation zur Kreditwürdigkeit oder das Sprachmodell hinter ChatGPT – aus Daten lernen und daraus sinnvolle Ergebnisse ableiten kann?

Unabhängig davon, ob es sich um eine einfache Klassifikation, eine Regression oder eine generative Anwendung handelt, müssen beim Aufbau eines KI-Modells insbesondere drei zentrale Bereiche adressiert werden:

- **Aufbereitung der Daten:** Wie müssen die vorhandenen Daten vorbereitet, bereinigt und strukturiert werden, damit das Modell relevante Muster erkennen kann?
- **Training und Auswahl des Modells:** Welche Daten werden für das Training oder die Anpassung verwendet, welches Modell eignet sich für den konkreten Anwendungsfall und wie wird dieses Modell konfiguriert?
- **Überprüfung und Bewertung des Modells:** Wie kann beurteilt werden, ob das Modell verlässliche Ergebnisse liefert und wie gut es auf neue, unbekannte Daten übertragbar ist? Wie kann Feedback wieder in das Modell zurückfließen?

Auch wenn jeder dieser Schritte im Detail unterschiedlich für verschiedene KI-Systeme ausgestaltet sein kann – etwa, wenn bei Sprachmodellen zusätzlich zwischen Pre-Training und Post-Training unterschieden wird –, bilden diese drei Bereiche den Kern des Modellierungsprozesses und werden in den folgenden Abschnitten Schritt für Schritt

5 Künstliche Intelligenz

erläutert. Dabei ist zu beachten, dass es sich nicht um einen streng linearen Ablauf handelt. Vielmehr ist die Modellentwicklung in der Praxis häufig iterativ: Erkenntnisse aus späteren Phasen führen dazu, dass frühere Schritte angepasst oder wiederholt werden müssen.

Data Pre-Processing

Bevor ein Modell überhaupt trainiert werden kann, müssen die Daten so vorbereitet werden, dass das Modell darin die relevanten Muster erkennen kann. Das Prinzip ist simpel: Garbage in, Garbage out. Wenn die Daten unvollständig, widersprüchlich oder schlecht strukturiert sind, wird auch ein noch so gutes Modell keine verlässlichen Resultate liefern. Dies gilt für jegliche Arten der AI, doch es gibt Möglichkeiten die Datenqualität durch sogenanntes Data Pre-Processing (Datenvorverarbeitung) aufzubessern:

- **Fehlende Werte (Imputing):** Modelle können mit leeren Feldern meist nicht direkt umgehen. Fehlende Angaben können je nach Kontext durch einfache Werte (z. B. Durchschnitt), durch eine eigene Kategorie (z. B. «unbekannt» oder durch eine Schätzung aus anderen Variablen ersetzt werden. Spannend ist: Manchmal wird dafür bereits ein Modell genutzt, um Daten für ein anderes Modell «zu reparieren»
- **Skalierung und Normalisierung:** Viele Verfahren reagieren empfindlich darauf, wenn Variablen sehr unterschiedliche Wertebereiche haben (z. B. Alter 18–100 versus Vermögen 0–1000 000). Darum werden numerische Spalten oft so transformiert, dass sie vergleichbarer werden (z. B. Mittelwert 0, Standardabweichung 1). Das kann Training und Ergebnisqualität deutlich verbessern.
- **Aussenseiter (Outliers) und Plausibilisierung:** Extremwerte oder offensichtliche Fehler (z. B. «Alter = 999») verfälschen Muster. Je nach Fall werden solche Werte bereinigt, gekappt oder gesondert behandelt.
- **Feature Engineering und Datenanreicherung:** Häufig entstehen die aussagekräftigsten Merkmale erst durch Umformungen. Im Control-ling wären das z. B. Zahlungsziel in Tagen, DSO, Abweichungen zum Budget, rollierende 3 Monats

Mittelwerte, Saisonindikatoren oder die Anzahl Mahnstufen. Solche Features machen Muster oft erst sichtbar.

- **Umgang mit unbalancierten Daten (Oversampling):** Wenn ein Ereignis selten ist (z. B. Zahlungsausfall oder Betrugsfall), könnte das Modell «lernen», fast immer den häufigen Fall vorherzusagen. Techniken wie Oversampling helfen, das Verhältnis im Training auszugleichen.

Bei generativen Sprachmodellen ist die Datenaufbereitung ebenfalls zentral, unterscheidet sich jedoch deutlich von der klassischen Aufbereitung tabellarischer Daten. Vor dem eigentlichen Pre-Training wird ein sehr grosser Textkorpus so vorbereitet, dass das Modell robuste und verallgemeinerbare sprachliche Muster lernen kann. Typische Schritte sind:

- **Quality filtering:** Entfernen von Spam, sehr kurzen oder inhaltlich schwachen Texten sowie automatisch generierten oder offensichtlich fehlerhaften Inhalten (z. B. reiner Boilerplate-Text von Webseiten).
- **Deduplication:** Entfernen von Duplikaten und nahe identischen Texten, damit das Modell nicht überproportional häufig dieselben Inhalte sieht und die Trainingsdaten nicht verzerrt werden.
- **Normalization & cleaning:** Vereinheitlichung von Zeichencodierung und Format, Entfernen technischer Artefakte (z. B. HTML-Reste) sowie konsistente Behandlung von Sonderzeichen und Leerzeichen.
- **Language & domain detection:** Erkennung von Sprache und Texttyp (z. B. Fliesstext, Code, Tabellenfragmente), um die gewünschte Mischung der Trainingsdaten zu steuern und unerwünschte Inhalte auszufiltern.
- **Sensitive content filtering:** Entfernen oder Maskieren von personenbezogenen Daten sowie anderen sensiblen oder regulatorisch problematischen Inhalten, um Datenschutz- und Compliance-Risiken zu reduzieren.
- **Segmentation & tokenization:** Aufteilung langer Dokumente in verarbeitbare Abschnitte (Chunking) und Übersetzung des Textes in Tokens, die das Modell im Training verarbeitet. Dabei

5 Künstliche Intelligenz

wird festgelegt, wie mit sehr langen Dokumenten umgegangen wird (z. B. Overlap zwischen Chunks).

Wichtige Begriffe: Pre-Processing, Imputing, Normalization, Feature Engineering, Outlier, Oversampling, Datenanreicherung.

Modellauswahl und Training

Sobald die Daten bereit sind, kann das Modell trainiert werden. Ziel des Trainings ist es, aus den vorhandenen Daten Muster zu erkennen und diese Muster im Modell abzubilden, sodass sie später auf neue, bisher unbekannte Daten angewendet werden können. Die Mustererkennung findet dabei innerhalb des Modells statt: Durch das Training wird die Struktur des Modells so angepasst, dass die in den Daten enthaltenen Zusammenhänge bestmöglich repräsentiert sind.

Die Wahl des Modells

Bevor ein Modell trainiert wird, muss jedoch entschieden werden, welches Modell für die jeweilige Fragestellung geeignet ist. Für die verschiedenen Arten der künstlichen Intelligenz stehen unterschiedliche Modelle zur Verfügung. Nachfolgend eine Auswahl von Modellen, die man kennen sollte:

- **Klassifizierung:** Decision Tree, Random Forest, XGBoost, Logistic Regression, GLM, Neural Networks
- **Regression:** Linear Regression, Ridge Regression, Lasso Regression, Polynomial Regression, Bayesian Linear Regression, Elastic Net Regression
- **Clustering:** K-Means, Nearest Neighbor

Die einzelnen Modelle im Detail zu verstehen, würde den Rahmen dieses Skripts sprengen. Stattdessen fokussieren wir uns bewusst auf ein einzelnes Modell, um zu verdeutlichen, was künstliche Intelligenz bzw. Mustererkennung im Kern bedeutet. Als Referenzmodell verwenden wir den **Decision Tree**, da sich dessen Funktionsweise besonders gut nachvollziehen lässt.

Training erklärt am Beispiel Decision Tree

Ein Decision Tree (Entscheidungsbaum) kann als eine Abfolge von Fragen an die Daten verstanden werden, die jeweils mit Ja oder Nein beantwortet werden. Abhängig von der Reihenfolge und den Antworten auf diese Fragen wird ein Datenpunkt schrittweise weitergeleitet, bis am Ende eine Klassenzuordnung oder Vorhersage resultiert.

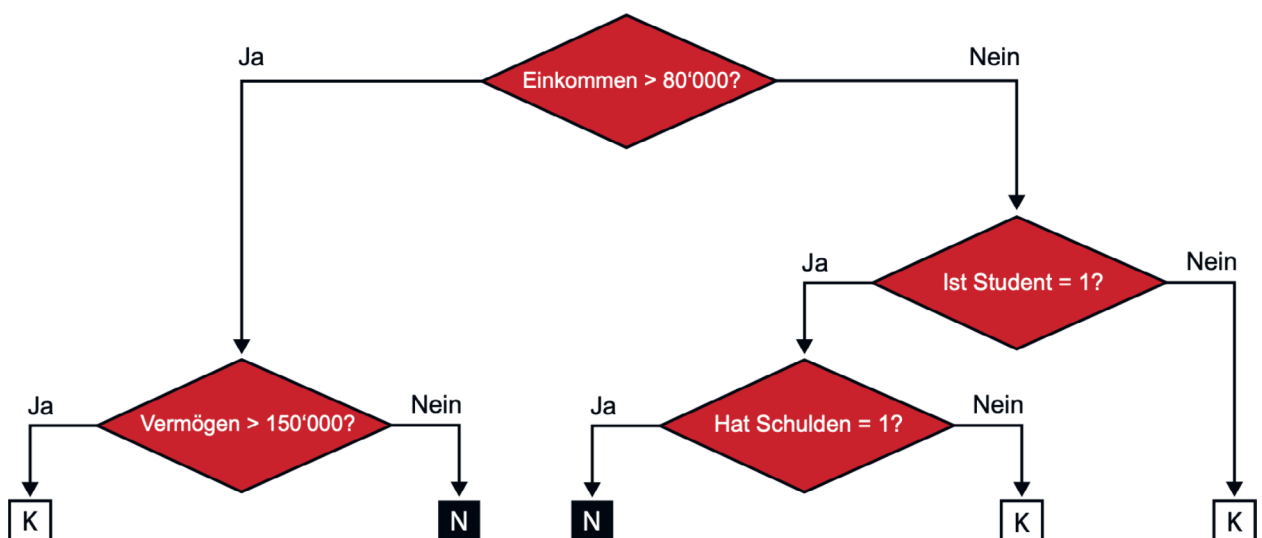


Abbildung 41: Ein einfacher beispielhafter Decision Tree zur Entscheidung auf Kreditwürdigkeit (eigene Darstellung)

5 Künstliche Intelligenz

Abbildung 41 illustriert dieses Prinzip am Beispiel der Kreditwürdigkeit. Nehmen wir an, eine Person ist Student, verfügt über ein jährliches Einkommen von CHF 30 000 und hat keine Schulden. Der erste Entscheidungsknoten prüft, ob das Einkommen höher als CHF 80 000 ist. Da dies nicht der Fall ist, wird die Person zum nächsten Knoten weitergeleitet, der überprüft, ob es sich um einen Studenten handelt. Diese Frage wird mit Ja beantwortet, woraufhin im nächsten Schritt geprüft wird, ob Schulden bestehen. Da keine Schulden vorliegen, resultiert am Ende die Einstufung als kreditwürdig (K).

Das Prinzip ist einfach – doch wo ist hier die Intelligenz? Um dies zu beantworten, stellen wir uns ein paar Fragen:

- Warum wird Einkommen zuerst abgefragt und nicht etwa, ob Schulden vorhanden sind?
- Weshalb liegt der Grenzwert bei CHF 80 000 und nicht beispielsweise bei CHF 50 000?
- Warum wird in der nächsten Verzweigung entweder nach Vermögen oder nach dem Studentenstatus gefragt?

Genau hier kommt KI ins Spiel. Die Fragen, ihre jeweiligen Grenzwerte sowie die Reihenfolge, in der sie gestellt werden, wurden im Training auf Basis historischer Daten ermittelt. Ziel dieses Trainings ist es, die in den Daten enthaltenen Muster so im Modell abzubilden, dass für einen möglichst grossen Teil der Trainingsdaten die korrekte Vorhersage resultiert.

Mit der Wahl eines Modells wird zunächst nur der **Rahmen der möglichen Entscheidungslogik** festgelegt. Erst durch das Training, also durch das systematische Erkennen von Mustern in den Daten, wird dieser Rahmen konkret ausgefüllt und für die Vorhersage oder Einstufung neuer Daten nutzbar gemacht.

Konfiguration des Modells vor dem Training

Ein Modell ist jedoch nicht gleich ein Modell – bereits **vor dem Training** müssen bestimmte Konfigurationsentscheidungen getroffen werden. Wenn ein Modell im Training gezielt auf die Trainingsdaten optimiert wird, stellt sich z. B. die Frage, weshalb die Komplexität des Modells begrenzt werden muss. Könnte ein Decision Tree nicht beliebig tief wachsen (also so viele Fragen, wie es will fragen), um die Trainingsdaten perfekt abzubilden?

Die Antwort ergibt sich, wenn wir den Blick über die Trainingsdaten hinaus richten. Wird ein Modell ausschliesslich auf die Trainingsdaten zugeschnitten, kann es diese zwar perfekt abbilden, verliert jedoch seine Fähigkeit zur Verallgemeinerung. Dieses Phänomen wird als **Overfitting** bezeichnet. Ein überangepasstes Modell liefert für die Trainingsdaten sehr gute Ergebnisse, scheitert jedoch häufig bei neuen, bisher unbekannten Daten.

Um Overfitting zu vermeiden und das Modell darauf zu beschränken, nur die wesentlichen Muster aus den Trainingsdaten zu übernehmen, wird die **Komplexität des Modells** begrenzt. Im Fall eines Decision Trees erfolgt dies z. B. durch die Einschränkung der maximalen Baumtiefe.

Die maximale Tiefe eines Decision Trees ist ein **konfigurierbarer Parameter**. Wird dieser zu niedrig gewählt, kann das Modell die zugrunde liegende Logik des Problems nicht ausreichend abbilden. Wird er zu hoch gewählt, steigt das Risiko von Overfitting. Neben der Baumtiefe existieren weitere Parameter, die die Struktur und das Verhalten eines Modells beeinflussen. Die systematische Suche nach einer geeigneten Kombination dieser Parameter wird als **Hyperparameter Optimization** bezeichnet.

Wichtige Begriffe: Decision Tree (Entscheidungsbaum), Modellparameter, Baumtiefe, Overfitting, Hyperparameter Optimierung.

5 Künstliche Intelligenz

Modell-Evaluierung

Ein zentraler Punkt dieses Abschnitts ist bislang noch offengeblieben: **Wann ist ein Modell gut?** Und wie lässt sich beurteilen, ob ein Modell besser ist als ein anderes?

Um diese Fragen zu beantworten, müssen zunächst geeignete **Testdaten** definiert werden. Testdaten sind Daten, bei denen das tatsächliche Ergebnis bekannt ist und die dazu verwendet werden, die Vorhersagen des Modells zu überprüfen. Da das Modell bereits auf den Trainingsdaten gelernt hat, ist es ungeeignet, diese Daten auch zur Bewertung heranzuziehen.

Eine **Metrik** beschreibt das Mass, anhand dessen die Qualität eines Modells beurteilt wird. In diesem Abschnitt beschränken wir uns auf Metriken zur Bewertung **binärer Klassifikatoren**, also Modelle mit zwei möglichen Ausprägungen des Ergebnisses.

Eine naheliegende Metrik zur Bewertung eines binären Klassifikationsmodells ist die **Accuracy**. Sie beschreibt den Anteil der korrekt klassifizierten Datenpunkte an allen betrachteten Datenpunkten.

Die **Confusion Matrix** (Abbildung 42) stellt die vom Modell vorhergesagten Klassen den tatsächlich eingetretenen Klassen gegenüber und bildet die Grundlage für die quantitative Bewertung eines Klassifikators:

- **True Positive (TP):** Das Modell sagt die relevante Klasse korrekt voraus (Kunde wurde als kreditwürdig vorhergesehen und ist es auch)
- **True Negative (TN):** Das Modell sagt die nicht relevante Klasse korrekt voraus (Kunde wurde als nicht kreditwürdig angesehen und ist es auch)
- **False Positive (FP):** Das Modell sagt fälschlicherweise die relevante Klasse voraus. (Kunde wurde als kreditwürdig vorhergesehen, ist es aber nicht)
- **False Negative (FN):** Das Modell erkennt die relevante Klasse fälschlicherweise nicht (Kunde wurde als nicht kreditwürdig vorhergesehen, obwohl er es war)

		Realität	
		Positive	Negative
Einschätzung durch KI	Positive	True Positive TP	False Positive FP
	Negative	False Negative FN	True Negative TN

Abbildung 42: Confusion Matrix (eigene Darstellung)

Optimalerweise befinden sich so viele Datenpunkte wie möglich bei TP und TN und so wenige wie möglich bei FN und FP. Die Accuracy berechnet sich als Anteil der korrekt klassifizierten Datenpunkte an allen Datenpunkten:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}$$

Die Accuracy ist jedoch nicht in allen Situationen eine geeignete Metrik. Dies gilt insbesondere bei **unausgeglichene Klassenverteilungen**. Bestehen beispielsweise 99 % der Datenpunkte aus der Klasse «nicht kreditwürdig», erzielt ein Modell, das **konsequent nur diese eine Klasse vorhersagt**, eine Accuracy von 99 %, ist für die Praxis jedoch wenig nützlich.

Um dieses Problem zu adressieren, wird häufig die Metrik **Precision** verwendet. Die Precision beschreibt den Anteil der korrekt als positiv klassifizierten Datenpunkte an allen Datenpunkten, die vom Modell als positiv vorhergesagt wurden:

$$\text{Precision} = \frac{TP}{TP + FP} = \text{«Wie oft haben wir richtig als Positive markiert?»}$$

5 Künstliche Intelligenz

Eine weitere zentrale Metrik ist der **Recall**. Der Recall beschreibt den Anteil der tatsächlich positiven Datenpunkte, die vom Modell korrekt als positiv erkannt werden:

$$\text{Recall} = \frac{TP}{TP + FN} = \text{«Wie viele Positives haben wir erkennen können?»}$$

Precision und Recall stehen häufig in einem Zielkonflikt. Wird die Precision erhöht, sinkt oft der Recall, und umgekehrt. Welche der beiden Metriken wichtiger ist, hängt davon ab, welche Art von Fehlentscheidung schwerer wiegt.

Um Precision und Recall gemeinsam zu berücksichtigen, wird häufig der **F1-Score** verwendet. Der F1-Score ist das harmonische Mittel von Precision und Recall:

$$F1 = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$$

Metrik	Bedeutung	Berechnung	Wann besonders wichtig
Accuracy	Anteil aller korrekt klassifizierten Fälle	$\frac{TP + TN}{TP + TN + FP + FN}$	Bei ausgeglichenen Klassen
Precision	Anteil korrekt positiver Vorhersagen	$\frac{TP}{TP + FP}$	Wenn False Positives teuer sind
Recall	Anteil korrekt erkannt positiver Fälle	$\frac{TP}{TP + FN}$	Wenn False Negatives teuer sind
F1-Score	Kompromiss zwischen Precision und Recall	$2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}}$	Wenn beide gleichermassen wichtig sind

Abbildung 43: Zusammenfassung der wichtigsten Metriken inkl. der jeweiligen Relevanz (eigene Darstellung)

Die präsentierten Metriken beziehen sich auf Klassifizierungsmodelle. Für Regressions- sowie Clusteringmodelle kommen andere Metriken zum Einsatz, da hier Accuracy, Precision, Recall und F1-Score keinen Sinn ergeben.

Wichtige Begriffe: Testdaten, Metrik, Confusion Matrix, True Positive (TP), True Negative (TN), False Positive (FP), False Negative (FN), Accuracy, Precision, Recall, F1-Score.

5.4 LLMs

Im vorhergehenden Kapitel haben wir gesehen, dass das Training eines Modells darauf abzielt, **Muster in Daten zu erkennen und diese Muster im Modell abzubilden**, um Vorhersagen für neue, bisher unbekannte Daten zu ermöglichen. Dieses Grundprinzip gilt nicht nur für numerische Daten oder Klassifikationsprobleme, sondern auch für Sprache.

Mustererkennung bei LLMs

Sprachmodelle (sogenannte large language models oder LLMs) folgen einem einfachen Ansatz: **Sie sagen fortlaufend das nächste Wort voraus**, basierend auf den zuvor gesehenen Wörtern. Das Trainingsziel besteht darin zu lernen, **welche Wörter mit welcher Wahrscheinlichkeit auf eine bestimmte Wortfolge folgen**. Auch hier handelt es sich nicht um ein inhaltliches Verständnis im menschlichen Sinne, sondern um **Mustererkennung in den Trainingsdaten**.

Während beim Decision Tree gelernt wurde, welche Fragen in welcher Reihenfolge mit welchen Grenzwerten sinnvoll sind, lernen LLMs, **welche sprachlichen Fortsetzungen typischerweise auf Fragen, Aussagen oder Satzanfänge folgen**. Diese statistischen sprachlichen Zusammenhänge werden im Training aus grossen Textmengen abgeleitet.

Um noch weiter zu präzisieren, werden eigentlich nicht Wörter, sondern sogenannte Tokens vorhergesagt. Tokens können grob als Teile von Wörtern oder ganze Wörter/Satzteile verstanden werden. Sie sind das «Verständigungsmittel» von LLMs.

5 Künstliche Intelligenz

Beispielsweise könnte der Satz «Die Abschreibungen sinken» durch die Tokens «Die» | «Abschreib» | «ungen» | «sink» | «en» | «.» dargestellt werden.

Die meisten Nutzer von GenAI stossen spätestens bei der Rechnungsstellung auf den Begriff Token: So wird die **Nutzung der Modelle** von Anbietern wie OpenAI, Mistral, Claude oder anderen **zumeist in Tokens dargestellt**. Stand Mai 2026 kosten 1 Mio. sog. input tokens (die in das Modell eingespeist werden), \$5.00. Auch *output token* werden berechnet, also die Länge der Ausgabe des Modells (derzeit bei \$30 pro 1 Mio. *output tokens*³⁴.

Die stochastische Natur von Sprachmodellen

Ein wesentliches Merkmal von Sprachmodellen ist ihre **stochastische Arbeitsweise**:

1. Für eine gegebene Token-Sequenz existieren in der Regel mehrere plausible nächste Tokens, denen Wahrscheinlichkeiten zugeordnet sind, die zusammen 100% ergeben.

2. Das Modell wählt anschliessend eines dieser Tokens aus – **nicht deterministisch**, sondern so, dass die Auswahl der jeweiligen Wahrscheinlichkeit entspricht.
3. Daraus folgt, dass bei identischer Eingabe **nicht zwingend dieselbe Ausgabe erzeugt wird**. Die Variation entsteht aus der Auswahl zwischen mehreren wahrscheinlichen Fortsetzungen und ist ein inhärenter Bestandteil des Modells.

Obwohl die eingeschränkte Reproduzierbarkeit der Ergebnisse **eine der grössten Herausforderungen beim Einsatz von Sprachmodellen im Unternehmensumfeld** darstellt, ist sie zugleich eine **wesentliche Voraussetzung für deren Leistungsfähigkeit**. Sie ermöglicht es Sprachmodellen, flexibel auf unterschiedliche Kontexte zu reagieren und sprachlich natürliche, nicht schematische Texte zu erzeugen.

Bestehender Text	Vorhersage des nächsten Wortes (inkl. Wahrscheinlichkeit)	
Die Rechnung ist noch nicht bezahlt, daher wird sie als	Forderung	46%
	offen	22%
	Aussenstand	14%
	Verbindlichkeit	10%
	Mahnung	8%

Abbildung 44: Nächste-Wort-Vorhersage eines Sprachmodells (der Einfachheit nicht als Nächste-Token-Vorhersage dargestellt). Man kann sich die Wahl des nächsten Tokens als «Ratespiel» vorstellen. Dabei rät das Sprachmodell zufällig eine Zahl zwischen 0 und 100. Wenn die Zahl kleiner als 46 ist, wird «Forderung» als nächstes Wort gewählt. Ist die Zahl zwischen 46 und 68 (46+22), dann wird das Wort «offen» gewählt, etc.

³⁴<https://openai.com/api/pricing/>

5 Künstliche Intelligenz

Metriken und Benchmark

Beim Training von Sprachmodellen kommen Metriken zum Einsatz, die messen, **wie gut ein Modell sprachliche Muster vorhersagen kann**. Zentrale Grössen sind hierbei zwar z. B. Perplexity oder BLEU Score, allerdings erfahren diese wenig Prominenz in der Kommunikation rund um LLMs. Stattdessen werden in der Regel Ergebnisse auf **standardisierten Benchmarks** präsentiert.

Benchmarks messen die Leistungsfähigkeit von LLMs bei **konkreten Aufgaben**, etwa beim logischen Schlussfolgern, Textverständnis oder Problemlösen. Sie liefern eine anwendungsnähere Einschätzung der Modellleistung. In Modellankündigungen wird häufig gezeigt, wie ein neues Modell auf diesen Benchmarks **im Vergleich zu früheren Versionen oder konkurrierenden Modellen** abschneidet.

Benchmark	Was wird gemessen	Typischer Score
MMLU / MMLU-Pro	Allgemeines Wissens- und Schlussfolgerungsvermögen	Prozent korrekt
GSM-8K	Mehrstufige mathematische Problemlösung	Prozent korrekt
HumanEval	Programmier- und Code-Generierungsfähigkeit	Prozent bestandene Aufgaben
Reading Comprehension (QA)	Textverständnis und Beantwortung komplexer Fragen	Prozent korrekt
Multi-Task Benchmarks	Leistung über verschiedene Aufgabentypen hinweg	Aggregierter Score

Reasoning

In den vorhergehenden Abschnitten wurde gezeigt, dass Sprachmodelle auf der Vorhersage des nächsten Tokens basieren und dabei Muster in Sprache erkennen. In vielen Fällen beschränkt sich diese Mustererkennung auf einfache Fortsetzungen oder Wissensabfragen. In anderen Fällen ist jedoch eine **mehrschrittige Herleitung** erforderlich, die man Reasoning nennt.

Beim Reasoning erzeugt ein Sprachmodell eine Abfolge von argumentativen Zwischenschritten, die zu einer Antwort hinführen. Diese Vorgehensweise wird häufig als **Chain-of-Thought (CoT)** bezeichnet. Dabei handelt es sich nicht um logisches Denken im menschlichen Sinne, sondern um sprachliche Muster, die im Training häufig gemeinsam auftreten und deshalb vom Modell reproduziert werden.

Der Unterschied lässt sich gut anhand von Beispielen aus dem Rechnungswesen verdeutlichen:

- **Kein Reasoning erforderlich:**

«Wie hoch ist der aktuelle Mehrwertsteuersatz in der Schweiz?»

→ einfache Wissensabfrage

- **Reasoning erforderlich:**

«Ist ein Sachverhalt als Rückstellung oder als Eventualverbindlichkeit zu behandeln»

→ Abwägung mehrerer Kriterien, Interpretation und strukturierte Argumentation

Sprachmodelle können solche Argumentationsketten sprachlich konsistent darstellen. Die einzelnen Zwischenschritte werden jedoch **nicht auf fachliche Korrektheit geprüft**, sondern lediglich als plausible Fortsetzung erzeugt.

Die Relevanz von Kontext

Sprachassistenten wie ChatGPT generieren Antworten ausschliesslich auf Basis der Informationen, die ihnen im jeweiligen Moment zur Verfügung stehen. Ohne zusätzlichen **Kontext** greifen sie auf allgemeines Wissen zurück, das im Training des Modells erlernt wurde. Dieses Wissen ist zwangsläufig unvollständig, teilweise veraltet (Stichwort Cut-off Date) und nicht unternehmensspezifisch. Entsprechend bleiben Antworten ohne Kontext oft generisch oder für den praktischen Einsatz im Rechnungswesen nur eingeschränkt brauchbar.

Gerade im Accounting ist Kontext zentral: Fragestellungen beziehen sich häufig auf interne Richtlinien, konkrete Verträge, unternehmensspezifische Kontierungslogiken, aktuelle Zahlen oder spezifische regulatorische Auslegungen. Ohne Zugriff auf

5 Künstliche Intelligenz

diese Informationen kann auch ein leistungsfähiger Sprachassistent wie ChatGPT keine belastbaren Aussagen treffen.

Grundsätzlich lassen sich drei Formen von Kontext unterscheiden, über die ChatGPT und vergleichbare Systeme heute verfügen können.

Kontext durch Prompt (expliziter Kontext)

Die einfachste Form von Kontext wird direkt im Prompt bereitgestellt. Der Nutzer ergänzt seine Frage um relevante Hintergrundinformationen, Annahmen oder Einschränkungen.

Beispiel:

«Beurteile, ob ein Sachverhalt als Rückstellung zu erfassen ist. Gegeben ist ein laufendes Gerichtsverfahren mit einer Eintrittswahrscheinlichkeit von ca. 60 % und einer erwarteten Belastung von CHF 500 000 gemäss interner Risikobewertung.»

Je präziser und relevanter dieser Kontext ist, desto zielgerichteter fällt die Antwort von ChatGPT aus. In der Praxis zeigt sich: Nicht das formale Layout eines Prompts («Prompt Engineering») ist entscheidend, sondern die Qualität, Vollständigkeit und fachliche Relevanz der bereitgestellten Informationen. Der Kontext dominiert klar über Stilfragen.

Kontext durch Dokumente (Retrieval-Augmented Generation)

In vielen Anwendungsfällen ist es nicht praktikabel, den gesamten relevanten Kontext manuell in einen Prompt zu integrieren. Moderne Sprachassistenten wie ChatGPT können daher zusätzlich mit Dokumenten arbeiten. Dieses Vorgehen wird als *Retrieval-Augmented Generation (RAG)* bezeichnet.

Dabei werden Dokumente – etwa Accounting Manuals, Verträge, Abschlussunterlagen oder Prozessbeschreibungen – so aufbereitet, dass bei einer Anfrage automatisch die relevantesten Textstellen identifiziert und dem Modell als zusätzlicher Kontext zur Verfügung gestellt werden. ChatGPT generiert seine Antwort dann nicht auf Basis des allgemeinen Trainingswissens, sondern unter explizitem Bezug auf diese Inhalte.

Typische Anwendungsfälle im Rechnungswesen sind:

- Beantwortung von Fachfragen unter Bezug auf interne Richtlinien oder Weisungen
- Zusammenfassung und Einordnung von Abweichungen auf Basis konkreter Reporting-Dokumente
- Unterstützung bei Prüfungsfragen unter Einbezug vergangener Abschlüsse oder interner Dokumentation

RAG reduziert das Risiko von Halluzinationen deutlich, da Aussagen auf konkret bereitgestützte Quellen zurückgeführt werden können.

Kontext durch Systemanbindung

Die umfassendste Form von Kontext entsteht durch die direkte Anbindung von Systemen. ChatGPT kann dabei nicht nur auf statische Dokumente, sondern auf aktuelle Daten aus angebundenen Anwendungen zugreifen, etwa aus ERP-, BI- oder Dokumentenmanagement-Systemen.

Dadurch werden dynamische, operative Anwendungsfälle möglich, zum Beispiel:

- Beantwortung von Fragen zu aktuellen Buchungsständen oder Kennzahlen
- Analyse von Abweichungen auf Basis tagesaktueller Transaktionsdaten
- Unterstützung bei Abstimmungs-, Freigabe- oder Klärungsprozessen

In ChatGPT erfolgt dies heute über sogenannte Apps oder Systemanbindungen, die es erlauben, externe Datenquellen kontrolliert einzubinden. Der Zugriff ist dabei durch klare Regeln (Guardrails) begrenzt, um Datenschutz-, Sicherheits- und Compliance-Anforderungen zu erfüllen.

Einordnung

Kontext ist der zentrale Hebel, um ChatGPT von einem allgemeinen Sprachassistenten zu einem fachlich nutzbaren Werkzeug im Rechnungswesen zu machen. **Ohne Kontext** liefert das Modell plausible, aber häufig **generische Antworten**. **Mit aus-**

5 Künstliche Intelligenz

reichend hochwertigem, aktuellem und kontrolliertem Kontext kann ChatGPT hingegen **gezielt, konsistent und nachvollziehbar unterstützen**.

Für den praktischen Einsatz im Accounting gilt daher: Der Mehrwert entsteht nicht nur durch das Modell, sondern vielmehr durch die Qualität, Aktualität und Governance des bereitgestellten Kontexts.

Agenten

ChatGPT ist als generalistischer Sprachassistent konzipiert. Er kann eine breite Palette an Fragen beantworten und Aufgaben unterstützen, ohne auf einen einzelnen Anwendungsfall spezialisiert zu sein. Für viele produktive Einsatzszenarien im Rechnungswesen ist dieser generalistische Ansatz jedoch nur begrenzt ausreichend. Hier kommen sogenannte **Agenten** zum Einsatz.

Ein Agent ist ein **spezialisierte KI-Assistent**, der für eine klar definierte Aufgabe oder einen abgegrenzten Prozess konzipiert ist. Im Gegensatz zum allgemeinen ChatGPT-Assistenten folgt ein Agent festen Instruktionen, nutzt gezielt ausgewählte Werkzeuge und agiert innerhalb klar definierter Grenzen.

Zentrale Eigenschaften eines Agenten

Ein Agent lässt sich durch mehrere, klar abgrenzbare Attribute beschreiben:

1. Instruktionen (Rolle, Aufgabe, Vorgehen)

Instruktionen bilden den Kern eines Agenten und legen explizit fest:

- welche **Rolle** der Agent einnimmt (z. B. Accounting-Assistent, Abstimmungs-Agent, Prüfungs-Support)
- welche **Aufgaben** er ausführen soll
- welche **Vorgehensweise** einzuhalten ist (z. B. Prüfschritte, Prioritäten, Eskalationslogik)

Ziel: Konsistente, reproduzierbare Ergebnisse statt situativer Einzelantworten.

2. Modellwahl

Für jeden Agenten wird ein konkretes KI-Modell festgelegt:

- leistungsstarke Modelle für komplexe Analysen oder fachliche Argumentation
- kleinere, schnellere Modelle für einfache oder häufig wiederkehrende Aufgaben

Die Modellwahl ist ein Trade-off zwischen Qualität, Geschwindigkeit und Kosten.

3. Werkzeuge (Tools)

Agenten können über reine Textgenerierung hinaus aktiv handeln, z. B. durch:

- Zugriff auf Dokumente oder Wissensdatenbanken
- Ausführen von Berechnungen oder Code
- Abfragen von ERP-, BI- oder DMS-Systemen
- Erzeugen strukturierter Ausgaben (Tabellen, Checklisten, Reports)

Werkzeuge machen den Agenten prozessfähig und operativ nutzbar.

4. Eingabe- und Ausgabeformate

Ein Agent ist darauf ausgelegt, mit klar definierten Formaten zu arbeiten:

- erwartete Eingaben (z. B. Frage, Parameter, Dokument)
- strukturierte Ausgaben (z. B. Tabelle, Bewertung, Entscheidungsempfehlung)

Dies erleichtert die Integration in bestehende Arbeitsabläufe.

5. Guardrails (Rahmenbedingungen)

Guardrails begrenzen das Handeln des Agenten explizit:

- welche Daten verwendet werden dürfen
- welche Aktionen erlaubt oder ausgeschlossen sind
- welche Themen oder Aussagen tabu sind

Im Rechnungswesen sind Guardrails zentral für:

- Compliance
- Datenschutz
- Nachvollziehbarkeit und Governance

5 Künstliche Intelligenz

Agenten in ChatGPT

In ChatGPT lassen sich Agenten heute beispielsweise in Form von Custom GPTs umsetzen. Diese ermöglichen es auch Nutzenden ohne Programmierkenntnisse, spezialisierte Assistenten zu erstellen, indem Instruktionen, Datenquellen und Verhaltensregeln konfiguriert werden. Beispiele im Rechnungswesen:

- Abschluss-Agent mit vordefinierter Monatsabschluss-Checkliste
- Richtlinien-Agent, der ausschliesslich auf Basis interner Accounting Manuals antwortet
- Abstimmungs-Agent zur strukturierten Analyse von Differenzen zwischen Neben- und Hauptbuch

Agenten-Workflows

Für komplexe Prozesse reicht ein einzelner Agent oft nicht aus. In solchen Fällen werden **Agenten-Workflows** eingesetzt, bei denen mehrere spezialisierte Agenten zusammenarbeiten. Typischer Ablauf:

- Triage-Agent klassifiziert die Anfrage,
- Weiterleitung an spezialisierte Agenten (z. B. Reporting-, Kontierungs- oder Dokumentations-Agent),
- Konsolidierung der Ergebnisse und strukturierte Rückgabe an den Nutzer.

Agenten-Workflows eignen sich besonders für Prozesse mit klaren Teilschritten, etwa Prüfungsanfragen, Freigabeprozesse oder standardisierte Analysen.

Evals

Während die Leistungsfähigkeit von Large Language Models häufig anhand standardisierter Benchmarks beurteilt wird, ist dies für den produktiven Einsatz von Agenten nicht ausreichend. Benchmarks messen allgemeine Modellfähigkeiten, sagen jedoch wenig darüber aus, ob ein Agent eine konkrete fachliche Aufgabe im Unternehmenskontext zuverlässig erfüllt.

Für Agenten werden daher anwendungsnahe Evaluationskriterien definiert, sogenannte Evals (kurz für Evaluation Metrics). Evals messen nicht die allgemeine Modellstärke, sondern die Qualität der Ergebnisse im Hinblick auf einen klar abgegrenzten Business-Use-Case.

Typische Formen von Evals sind beispielsweise:

- **Fachliche Korrektheit:** Entsprechen die Ergebnisse den relevanten Rechnungslegungsstandards, internen Richtlinien oder definierten Entscheidungsregeln?
- **Vollständigkeit:** Wurden alle geforderten Aspekte berücksichtigt (z. B. alle Prüfschritte einer Abschluss-Checkliste)?
- **Konsistenz:** Liefert der Agent bei vergleichbaren Eingaben reproduzierbare Ergebnisse?
- **Struktur und Format:** Entspricht die Ausgabe dem erwarteten Format (z. B. Tabelle, Entscheidungsempfehlung, Begründung)?
- **Fehlerarten:** Wie häufig treten fachlich kritische Fehler auf (z. B. falsche Klassifikation oder fehlerhafte Kontierung)?

Entscheidend ist, dass Evals nicht abstrakt, sondern **geschäftsfähig** definiert werden. Die Auswahl der richtigen Evals ist häufig bereits die halbe Miete, da sie zwingt, den Anwendungsfall präzise zu formulieren und klar festzulegen, was aus fachlicher Sicht ein «gutes Ergebnis» darstellt.

Neben der Definition der Evals müssen zudem konkrete Testfälle festgelegt werden, die diese Kriterien abdecken. Diese Testfälle sind regelmässig – idealerweise automatisiert – auszuführen, um zu überprüfen, wie gut der Agent seine Aufgabe weiterhin erfüllt. Dies ist insbesondere relevant, wenn:

- auf ein **neues Modell** gewechselt wird
- **Instruktionen, Prompts oder Guardrails** angepasst werden
- sich **Datenquellen oder Systemanbindungen** ändern

5 Künstliche Intelligenz

Durch die wiederholte Ausführung derselben Testfälle kann systematisch beurteilt werden, ob sich die Qualität verbessert, stabil bleibt oder unbeabsichtigte Verschlechterungen auftreten. Evals entwickeln sich damit von einer einmaligen Bewertung zu einem kontinuierlichen Instrument für Qualitätssicherung, Vergleichbarkeit und Governance von KI-basierten Agenten.

Einordnung

Accountants werden in der Regel **keine eigenen KI-Modelle entwickeln oder trainieren**. Der Aufbau und das Training von Modellen erfordert spezialisiertes technisches Know-how sowie erhebliche Daten- und Infrastrukturressourcen.

Sehr wohl werden Accountants jedoch **eigene Agenten konzipieren und aufbauen**. Dabei geht es nicht um das Erstellen neuer KI-Modelle, sondern um die **fachliche Konfiguration bestehender Modelle**:

- Definition von Aufgaben, Rollen und Instruktionen
- Festlegung von Prüf- und Entscheidungslogiken
- Einbindung relevanter Dokumente, Datenquellen und Systeme
- Setzen von Guardrails zur Einhaltung von Compliance und Governance

Der fachliche Mehrwert entsteht damit nicht auf Ebene des Modells, sondern auf Ebene des **Agenten-Designs**. Accountants bringen hier ihre Domänenexpertise ein und übersetzen fachliche Anforderungen in strukturierte, reproduzierbare Arbeitsabläufe, die durch KI unterstützt werden.

Wichtige Begriffe: Large Language Models (LLMs), Token, Input Tokens, Output Tokens, stochastische Arbeitsweise, nächste-Token-Vorhersage, Benchmarks, Reasoning, Chain-of-Thought (CoT), Agenten, Agenten-Workflows, Guardrails, Tools, Kontext, Evals, Testfälle.

5.5 Anfälligkeiten eines KI-Modells

Bei der Realisierung einer KI muss man sich jederzeit der **Grenzen eines KI-Modells** bewusst sein. Ein solches Modell ist grundsätzlich nur so gut wie die Daten, die ihm zur Verfügung gestellt werden, und wie diese Daten im Training verarbeitet werden. Weisen die Trainingsdaten eine unzureichende Qualität auf oder sind sie nicht in geeigneter Form aufbereitet, können selbst umfangreiche Anpassungen der Modellparameter kein zufriedenstellendes Ergebnis erzielen.

Darüber hinaus ergeben sich Anfälligkeiten nicht nur aus den Daten selbst, sondern auch aus der **Funktionsweise des jeweiligen Modells**. Insbesondere bei Sprachmodellen können dadurch systematische Effekte auftreten, die nicht unmittelbar auf einzelne Trainingsdaten zurückzuführen sind.

Vor diesem Hintergrund stellt sich die Frage, **welche typischen Anfälligkeiten bei KI-Modellen auftreten können**.

Overfitting

Overfitting beschreibt den Zustand, in dem sich ein Modell im Training zu stark an die Trainingsdaten anpasst. In diesem Fall ist die Modellleistung auf den Trainingsdaten sehr hoch, während die Qualität der Ergebnisse auf neuen, unbekannten Daten deutlich abnimmt. Das Modell hat dann nicht die zugrundeliegenden Muster des Problems erlernt, sondern spezifische Besonderheiten der Trainingsdaten übernommen.

Overfitting kann reduziert werden, indem die Modellkomplexität und die Parameter so gewählt werden, dass eine **hinreichende Verallgemeinerung** erzwungen wird. Dies gilt sowohl für klassische KI-Modelle als auch für generative Modelle, bei denen sich Overfitting beispielsweise in einer zu starken Orientierung an gesehenen Textmustern oder Formulierungen äussern kann.

5 Künstliche Intelligenz

Zur Identifikation geeigneter Parameterkonfigurationen wird häufig **Hyperparameter Optimization** eingesetzt, wobei die Modellleistung auf separaten Testdaten bewertet wird. Um zu vermeiden, dass sich ein Modell auch an diese Testdaten anpasst, kommt in der Praxis oft die **Cross-Validation** zum Einsatz. Dabei wird die Modellgüte – etwa gemessen anhand geeigneter Metriken – über mehrere Trainingsdurchläufe hinweg aggregiert beurteilt.

Bias

Bias – oder Voreingenommenheit – bezeichnet den Zustand, in dem ein KI-Modell bestimmte Attribute oder Gruppen systematisch **bevorzugt oder benachteiligt**. Dies liegt vor, wenn Datenpunkte mit einem bestimmten Merkmal häufiger einer bestimmten Klasse zugeordnet werden, obwohl diese Zuordnung im Einzelfall sachlich nicht gerechtfertigt ist.

Ein typisches Beispiel zeigt sich bei der Beurteilung der Kreditwürdigkeit. Wird ein Modell unter anderem mit Daten trainiert, in denen eine bestimmte Postleitzahl stark mit abgelehnten Kreditgesuchen korreliert, kann dieses Merkmal überbewertet werden. In der Folge kann eine Bewerbung trotz stabilem Einkommen oder vorhandenem Vermögen negativ beurteilt werden – allein aufgrund dieses einen Attributs.

Bias ist häufig das Resultat **unausgewogener Trainingsdaten**. Eine sorgfältige Auswahl und die Aufbereitung der Daten sind daher zentral. Zusätzlich kann es sinnvoll sein, bestimmte Merkmale wie Alter, Geschlecht oder Herkunft **bewusst nicht als Features zu verwenden**, sofern sie für die Fragestellung nicht zwingend erforderlich sind.

Intransparenz

KI kann komplexe Zusammenhänge aus historischen Daten lernen und für Vorhersagen nutzen. In vielen Anwendungsfällen reicht eine reine Vorhersage jedoch nicht aus. Insbesondere im Finanz- und Rechnungswesen müssen Entscheidungen **nachvollziehbar begründet** werden.

Viele KI-Modelle sind in ihrer Entscheidungslogik nur **eingeschränkt transparent**. Es ist oft nicht unmittelbar ersichtlich, weshalb ein Modell für einen konkreten Fall eine bestimmte Einschätzung trifft. Diese Intransparenz stellt eine wesentliche Anfälligkeit dar, da sie die Überprüfbarkeit und Akzeptanz von KI-gestützten Entscheidungen erschwert.

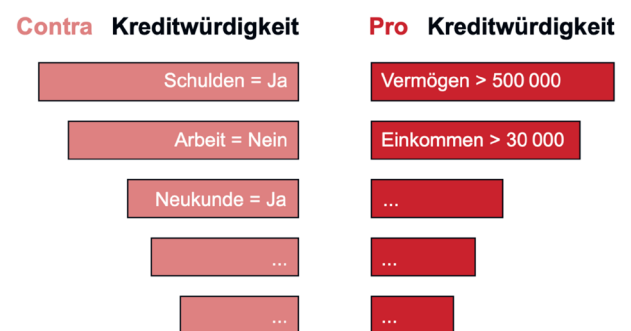


Abbildung 45: Explainable AI erklärt, welche Features pro und contra zur Kreditwürdigkeit stehen (eigene Darstellung)

Einen Ansatz zur Reduktion dieser Intransparenz bietet die sogenannte **Explainable AI (XAI)**. Ziel ist es, darzustellen, **welche Merkmale (Features)** für oder gegen eine Vorhersage gesprochen haben und wie stark deren Einfluss im konkreten Fall war.

Abbildung 45 zeigt dies am Beispiel der Kreditwürdigkeit. Links sind Merkmale dargestellt, die gegen eine Kreditwürdigkeit sprechen, rechts jene, die dafür sprechen. Die Länge der Balken spiegelt die relative Bedeutung der jeweiligen Merkmale wider. In diesem Beispiel überwiegen die Contra-Argumente.

Concept Drift

Die Muster, die ein KI-Modell aus historischen Daten lernt, sind nicht zwingend dauerhaft stabil. Beim Concept Drift verändern sich die Zusammenhänge zwischen Eingabedaten und Zielvariable im Zeitverlauf, wodurch ein Modell an Aussagekraft verlieren kann.

5 Künstliche Intelligenz

Am Beispiel der Kreditwürdigkeit erklärt: Wird ein Modell mit Daten aus einer Phase konservativer Kreditvergaberichtlinien trainiert, bildet es diesen Ansatz ab. Ändert sich die Strategie der Bank später, etwa durch den Verzicht auf bestimmte Kriterien, spiegelt das Modell diese Veränderung erst wider, wenn aktualisierte Trainingsdaten verwendet werden.

Concept Drift macht deutlich, dass KI-Modelle **regelmässig überwacht und bei Bedarf neu trainiert** werden müssen. Im Falle von GenAI muss meistens nicht das gesamte LLM neu trainiert werden (was für die meisten Unternehmen sowieso kaum möglich wäre), es reicht zumeist schon eine Aktualisierung des Prompts.

Halluzinationen

Eine besondere Anfälligkeit von Sprachmodellen sind sogenannte Halluzinationen. Darunter versteht man Aussagen, die inhaltlich falsch oder nicht belegbar sind, obwohl sie sprachlich korrekt und plausibel formuliert wirken.

Halluzinationen entstehen, weil Sprachmodelle darauf ausgelegt sind, **wahrscheinliche sprachliche Fortsetzungen** zu erzeugen. Sie prüfen Aussagen nicht auf inhaltliche Richtigkeit, sondern orientieren sich an Mustern aus den Trainingsdaten und dem aktuellen Kontext. Fehlen klare Informationen oder ist der Kontext unvollständig, kann das Modell plausibel, aber falsche Inhalte generieren.

Im Rechnungswesen ist diese Anfälligkeit besonders kritisch, da halluzinierte Zahlen, Regeln oder Interpretationen auf den ersten Blick überzeugend erscheinen können. Halluzinationen lassen sich nicht vollständig vermeiden, sondern nur durch geeignete Massnahmen reduzieren, etwa durch klar formulierte Anfragen, den Einbezug verlässlicher Datenquellen und eine konsequente fachliche Prüfung der Ergebnisse.

Wichtigste Begriffe: Bias (Voreingenommenheit), Intransparenz, Explainable AI (XAI), Concept Drift, Halluzinationen

5.6 Buy-vs.-Build

Sobald ein geeigneter Anwendungsfall identifiziert wurde, stellt sich im nächsten Schritt die Frage nach der Umsetzungsstrategie: Buy oder Build?

- **Buy:** Gibt es eine verfügbare Marktlösung, die den Grossteil der fachlichen und technischen Anforderungen abdeckt und mit vertretbarem Anpassungsaufwand eingesetzt werden kann?
- **Build:** Sind die Anforderungen des Unternehmens so spezifisch, dass keine geeignete Standardlösung existiert und eine Eigenentwicklung notwendig erscheint?

Im Rechnungswesen war eine **Build-Strategie** lange Zeit eher theoretischer Natur. Der Aufbau eigener KI-Lösungen war mit erheblichem Ressourcenaufwand verbunden und setzte spezialisiertes Know-how sowie entsprechende Infrastruktur voraus. Zudem haben viele Anbieter ihre bestehenden **ERP-, CRM- und Finanzsysteme** schrittweise um integrierte KI-Funktionalitäten erweitert, wodurch der Mehrwert eigenentwickelter Lösungen begrenzt war.

Beispiele integrierter KI-Funktionalitäten in Standardlösungen:

- **SAP (z. B. S/4HANA, SAP Analytics Cloud):** Automatisierte Belegerkennung, Buchungsvorschläge, Anomalieerkennung in Buchungen, Cashflow-Prognosen sowie KI-gestützte Analysen und Forecasts direkt auf ERP-Daten
- **Microsoft Dynamics / Business Central:** Unterstützung bei der Klassifikation von Transaktionen, Prognosen, Automatisierung von Abschlussprozessen und Integration mit Power BI und Copilot-Funktionalitäten
- **Oracle Financials / NetSuite:** KI-gestützte Forecasts, Abweichungsanalysen und Automatisierung von wiederkehrenden Finanzprozessen

Diese Lösungen verfolgen klar eine **Buy-Strategie**, bei der KI-Funktionalitäten eng in bestehende Systeme eingebettet sind und ohne Eigenentwicklung genutzt werden können.

5 Künstliche Intelligenz

Seit Oktober 2022 hat sich diese Ausgangslage jedoch spürbar verändert. Mit der Veröffentlichung von **ChatGPT** wurde erstmals ein allgemein verfügbarer KI-Assistent etabliert, der als **domänenübergreifendes Buy-Tool** menschliche Fragen in natürlicher Sprache beantworten und bei einer Vielzahl wissensintensiver Aufgaben unterstützen kann.

Beispiele für den Einsatz von ChatGPT im Rechnungswesen:

- Unterstützung bei der **Formulierung von Abschlusskommentaren** (z.B. Management Summary, Abweichungsanalysen)
- Zusammenfassung und sprachliche Aufbereitung von **Richtlinien, Rechnungslegungsstandards oder internen Weisungen**
- Hilfe bei der **Interpretation fachlicher Fragestellungen**, etwa zur Abgrenzung von Rückstellungen und Eventualverbindlichkeiten
- Unterstützung bei der **Vorbereitung von Präsentationen** für Geschäftsleitung oder Revision
- Strukturierung von **Checklisten** für Monats- oder Jahresabschlüsse

Noch weitergehend ermöglichen Werkzeuge wie ChatGPT inzwischen nicht nur die Nutzung eines vorgefertigten Assistenten, sondern auch den Aufbau **unternehmensspezifischer KI-Assistenten**. Über sogenannte **CustomGPTs** oder vergleichbare Konzepte lassen sich Assistenten gezielt konfigurieren und mit eigenen Datenquellen anreichern, etwa:

- Interne Richtlinien und Arbeitsanweisungen
- Prozessbeschreibungen des Rechnungswesens
- Interne FAQs oder Schulungsunterlagen

Damit entsteht ein **Hybridansatz zwischen Buy und Build**: Die zugrunde liegende KI-Technologie wird eingekauft, während Fachlogik, Kontext und Datenhoheit im Unternehmen verbleiben.

Wichtige Begriffe: Buy (Standardlösung), Build (Eigenentwicklung), Hybridansatz, Custom GPTs

5 Künstliche Intelligenz

Case Study: Ein Tag im Leben von Controllerin Lea, die vor kurzem KI für sich gefunden hat

08:30 – Tagesstart: Reporting, Excel und ein Abschreibungsmodell

Lea startet mit dem Monatsreporting. Sie nutzt ChatGPT nicht nur für Text, sondern als Arbeitsassistent für Modellierung:

Excel-Formelhilfe (konkret): Sie hat eine recht unschöne Pivot-/XLOOKUP-Kombination, die sie selbst nicht mehr versteht und lässt sich eine robuste Formel bauen (inkl. Fehlerbehandlung, z. B. IFERROR, sauberer Referenzen etc.).

Depreciation-Modell (konkret): Für neue Anlagen will sie ein Abschreibungsmodell (linear, optional degressiv, Startmonat/Teilperiodenlogik, Restwert, Anlagenabgänge). In ChatGPT erklärt sie, was sie braucht, wählt «Denkstufe» (engl. reasoning level) Thinking und bekommt nach ein paar Minuten Denkzeit ein schön formuliertes Excel-Worksheet mit mehreren Tabs. Lea muss zwar noch ein paar Anpassungen machen, ist aber glücklich, einen Vorschlag bekommen zu haben und nicht von null anfangen zu müssen.

Ergebnis: Stundenlanges Kreieren von Excel-Formeln und Debuggen ist jetzt in Minuten abgeschlossen.

10:30 – Ad-hoc: Ein Management-Update innerhalb von Minuten.

Die Chefin kommt mit einer spontanen Anfrage um die Ecke: «Kannst du die Abschreibungsmodellierung kurz in eine GL-taugliche Story bringen?»

Lea nutzt ChatGPT, um direkt Folieninhalte zu erzeugen. Sie verbindet ChatGPT mit ihrem Sharepoint, referenziert das Excel und fragt ChatGPT nach einer möglichen Struktur. Nach ein paar Iterationen ist sie zufrieden und lässt sich die Präsentation ausgeben.

12:00 – Deep Research: «Gab es so einen Fall schon mal?»

Bei einem grossen Kunden wurde ein mehrjähriger Servicevertrag geändert: Der Kunde bekommt ab Q3 einen rückwirkenden Rabatt und zusätzlich eine Sondergutschrift, die operativ als «Goodwill» verkauft wurde. Sie muss sich jetzt überlegen, wie sie das im ERP verbucht.

Lea weiss: Wenn sie das «aus dem Bauch» beantwortet, ist die Chance hoch, dass sie entweder einen anderen Weg als ihre Kollegen wählt oder im schlimmsten Fall sogar noch gegen das Gesetz verstösst. Also nutzt sie Deep Research, um bereits vorhandene interne Präzedenzfälle im eigenen System zu finden und ausserdem in externen Quellen zu suchen, worauf hier zu achten ist.

14:30 – Der Auslöser: Ständige ABAP-Fragen von Kollegen

Am Nachmittag fragt ein Kollege: «Wie ist eigentlich der ABAP-Prozess für die Zahlungszuordnung / Buchungslogik aufgebaut?»

Lea merkt: Diese Frage kommt nicht zum ersten Mal. Sie kann das jedes Mal erklären –oder sie baut einen Agenten, der das sauber und konsistent beantwortet (und am wichtigsten: ihr Arbeit abnimmt). Sie entscheidet sich für letzteres und baut einen Agenten, der den relevanten ABAP-Code und Prozessdoku kennt, ihre bewährten Antworten aus früheren Fragen als Wissensbasis nutzt, Antworten immer im gleichen Format liefert (Kurzantwort + Details + «wo im Code»):

5 Künstliche Intelligenz

Case Study: Ein Tag im Leben von Controllerin Lea, die vor Kurzem KI für sich gefunden hat

Instruktionen: «Beantworte Fragen nur mit Referenzen auf Code/Doku/Q&A. Wenn unsicher: frag nach Artefakt oder markiere Unsicherheit.»

Kontextquellen: ABAP-Repository/Export, Prozessbeschreibung, FAQ-Log (z.B. vergangene Ticket-Antworten).

Output-Format:
TL;DR (2–3 Sätze)

Prozessschritte (nummeriert)

Code-Referenz (File/Method/Comment)

Risiken/Edge Cases

Lea definiert 5–10 Standardfragen (Testfälle):

«Welche Bedingungen triggern Buchungstyp X?»

«Wo wird Feld Y gemappt?»

«Was passiert bei fehlender Referenz im Verwendungszweck?»

Auf diesen Standardfragen bewertet sie Qualität anhand einiger Evals, darunter:

Korrektheit + Referenzen: Kernaussagen müssen mit Code/Doku übereinstimmen und eine konkrete Referenz (File/Method/Comment) enthalten; sonst als unsicher markieren.

Vollständigkeit der Prozesslogik: Die Antwort deckt die relevanten Schritte und wichtigsten Abzweigungen/Edge Cases ab.

Format-Compliance: Ausgabe hält strikt das vereinbarte Format ein (TL;DR → Schritte → Code-Referenz → Risiken).

Lea ist zufrieden mit den Evals und teilt den Agenten mit allen Kollegen aus ihrem Team.

17:30 – Kurz vor Feierabend: Offene Posten sind noch nicht zugeordnet

Lea schaut in die Bankumsätze: Es gibt eingehende Überweisungen, die noch keinem offenen Debitorenposten zugeordnet wurden. Lea weiss aber, dass es in SAP eine Funktion gibt, die diese automatische Zuweisung unterstützt. Sie lässt die Zuordnung laufen und prüft anschliessend gezielt die Fälle, die unter 95 % Confidence liegen. Bei diesen Zahlungen öffnet sie die Erklär-Ansicht (Explainability/«Warum diese Zuordnung?») und versteht in wenigen Sekunden, auf welcher Logik die Zuordnung basiert. Einen Fall korrigiert sie manuell – aber insgesamt ist sie deutlich schneller durch. Dank der Unterstützung muss sie das Abendessen nicht allein im Büro verbringen.

Datensicherheit und Datenschutz



Ralph Hutter

6.1 Strategische Grundlagen: Daten als Finanzwert und Risiko

In der modernen Unternehmensführung hat sich ein Paradigmenwechsel vollzogen: Daten sind nicht mehr nur ein Nebenprodukt der Geschäftstätigkeit, das buchhalterisch erfasst und archiviert werden muss. Sie haben sich zu einem zentralen immateriellen Vermögenswert (Asset) entwickelt.

Gleichzeitig verschärfen die regulatorischen Vorgaben der Eidgenössischen Finanzmarktaufsicht (FINMA), das revidierte Datenschutzgesetz (DSG) sowie die Europäische Datenschutz-Grundverordnung (DSGVO) die Anforderungen an den operativen Umgang und den Schutz von Kundendaten massiv. Damit bergen Datenbestände und deren Bearbeitung auch Risiken (Liabilities), welchen aktiv begegnet werden muss.

Für das moderne Accounting ergibt sich daraus ein Spannungsfeld: Einerseits sollen Daten maximal genutzt und verknüpft werden, um Wert zu schöpfen (Business Intelligence) und andererseits gilt es, die zahlreichen regulatorischen Vorschriften zu berücksichtigen.

Für die Ausgestaltung der **Data Governance** ist eine präzise Unterscheidung zwischen Datenschutz (Data Privacy), Datensicherheit (Data Security) und Cyber-Sicherheit (Cyber Security) unerlässlich. Im alltäglichen Sprachgebrauch werden diese Begriffe allerdings oft synonym verwendet, doch für Accountants und Finanzverantwortliche implizieren sie unterschiedliche Pflichten und Massnahmen.

Einführung Datenschutz – Schutz der persönlichen Daten

Der Datenschutz bezieht sich auf den rechtlichen Rahmen und den Schutz der Persönlichkeitsrechte natürlicher Personen. Er regelt, ob, zu welchem Zweck und in welchem Umfang personenbezogene Daten verarbeitet werden dürfen.

- **Fokus:** Der Schutz des Individuums vor missbräuchlicher Datenverarbeitung.
- **Relevanz im DSG:** Das revidierte DSG verlangt Transparenz, Verhältnismässigkeit und Zweckbindung. Das bedeutet, dass Kundendaten nur für den definierten Zweck genutzt und nicht ohne Rechtsgrundlage für andere Analysen zweckentfremdet werden dürfen.
- **Ziel:** Einhaltung der Compliance-Vorgaben und Schutz der Privatsphäre.

Einführung Datensicherheit – Technische und organisatorische Massnahmen

Die Datensicherheit befasst sich mit den **technischen und organisatorischen Massnahmen** (TOMs), die notwendig sind, um Daten vor Bedrohungen zu schützen. Sie ist die technische Voraussetzung für effektiven Datenschutz.

- **Fokus:** Der Schutz der Datenbestände selbst. Die zentralen Schutzziele sind hierbei Vertraulichkeit (Zugriffsschutz), Integrität (Unverfälschtheit) und Verfügbarkeit (Ausfallsicherheit).
- **Relevanz im DSG:** Das Gesetz verpflichtet Unternehmen, eine angemessene Datensicherheit durch moderne Technologien zu gewährleisten («State of the Art»).
- **Ziel:** Verhinderung von Datenverlust, Diebstahl, Manipulation oder unbefugtem Zugriff.
- **Notfallplanung (Business Continuity):** Pläne zur Wiederherstellung der Buchhaltungssysteme

6 Datensicherheit und Datenschutz

nach einem Angriff, um die Geschäftstätigkeit schnellstmöglich wiederaufzunehmen und regelmässige Trainings dazu.

Einführung Cyber Security: Schutz vor externen Bedrohungen

Während sich die Datensicherheit auf den internen Schutzmechanismus der Daten konzentriert, befasst sich die **Cyber Security** mit der Verteidigung der gesamten Umgebung gegen Angriffe von aussen. Für Accountants ist dies von höchster Relevanz, da Finanzabteilungen aufgrund ihrer Zugriffsmöglichkeiten auf Zahlungsströme bevorzugte Ziele für Kriminelle sind.

Fokus: Abwehr von Angriffen auf Infrastruktur und Mensch

Der Fokus liegt auf dem aktiven Schutz der IT-Systeme, Netzwerke und Endgeräte gegen böswillige Akteure aus dem Cyberraum.

- **Technische Bedrohungen:** Abwehr von Malware, Ransomware (Verschlüsselungstrojaner) und DDoS-Attacken, die darauf abzielen, Systeme lahmzulegen.
- **Menschliche Angriffsvektoren:** Schutz vor Social Engineering (z. B. CEO-Fraud) und Phishing, bei denen Mitarbeitende im Accounting manipuliert werden, um Zahlungen auszulösen oder Zugangsdaten preiszugeben.

Relevanz im DSG

Sorgfaltspflicht und Meldewesen: Das DSG verlangt «angemessene technische und organisatorische Massnahmen» und definiert neu eine Meldepflicht.

Ziel: Aktiver Schutz der IT-Systeme, Netzwerke und Endgeräte sowie die Gewährleistung der betrieblichen Kontinuität.

Vergleich

Merkmal	Datenschutz (Data Privacy)	Datensicherheit (Data Security)	Cyber Security
Kernfrage	Dürfen wir diese Daten bearbeiten und wozu?	Wie schützen wir die Daten vor Verlust und Zugriff?	Wer greift uns an und wie wehren wir ab?
Schutzobjekt	Die natürliche Person (Persönlichkeitsrecht).	Der Datenbestand (das Asset selbst).	Die Infrastruktur & Umgebung (IT, OT, Cloud).
Hauptrisiko	Missbrauch, Gläserner Kunde, Rechtsverstoß.	Datenverlust, Manipulation, Hardware-Ausfall.	Hackerangriff, Ransomware, Sabotage.
Verantwortung	Legal, Compliance, DPO (Datenschutzberater).	IT-Leitung, Data Owner, CISO.	CISO, Security Operations Center (SOC).

Abbildung 46: Merkmale von Datenschutz, Datensicherheit und Cyber Security

Im Kontext des Schweizer Datenschutzgesetzes (DSG) sowie internationaler Standards (wie der DSGVO) sind Datenschutz, Datensicherheit und Cyber Security zwar eigenständige Disziplinen, sie bilden jedoch eine untrennbare **Symbiose**. Für das moderne Accounting ist es entscheidend, diese drei Gebiete nicht als isolierte Themen, sondern als integriertes Risikomanagement verstanden werden.

6.2 Datenschutz

Der Datenschutz bildet den rechtlichen Rahmen für den Umgang mit personenbezogenen Daten und fokussiert auf den Schutz der Persönlichkeit und der Grundrechte natürlicher Personen. Er regelt die Frage, ob und zu welchem Zweck Daten bearbeitet werden dürfen. Mit dem Inkrafttreten des revidierten Schweizer Datenschutzgesetzes haben sich die Anforderungen an Transparenz, Zweckbindung und Datensparsamkeit verschärft, was Unternehmen zu einer präzisen Dokumentation ihrer Bearbeitungstätigkeiten zwingt.

6 Datensicherheit und Datenschutz

Für die Unternehmensführung und das Finanzmanagement bedeutet dies, dass Daten nicht mehr unbegrenzt als «Rohstoff» verfügbar sind, sondern als reguliertes Gut betrachtet werden müssen. Die Einhaltung der Datenschutzvorgaben ist keine rein administrative Aufgabe, sondern eine Compliance-Anforderung, deren Verletzung direkte sanktionsrechtliche Konsequenzen für die verantwortlichen Leitungspersonen haben kann.

Grundprinzipien des revidierten Schweizer Datenschutzgesetzes

Das revidierte Gesetz behält die pragmatische Grundhaltung des Schweizer Rechts bei (oft als «Swiss Finish» bezeichnet) und nähert sich aber inhaltlich und terminologisch stark der EU-DSGVO an. Es basiert auf fundamentalen Prinzipien, die für jede Phase des Datenlebenszyklus gelten, von der Erhebung über die Speicherung bis zur Vernichtung.

Transparenz

Datenbearbeitungen müssen für die betroffene Person erkennbar und nachvollziehbar sein. Das bedeutet, dass Unternehmen proaktiv informieren müssen, welche Daten sie zu welchem Zweck sammeln. Verdeckte Profilbildungen oder die Nutzung von Daten zu Zwecken, die für den Betroffenen überraschend wären, sind unzulässig.

Verhältnismässigkeit (Datensparsamkeit)

Es darf nur das bearbeitet werden, was für den definierten Zweck absolut notwendig ist. Das Sammeln von Daten «auf Vorrat», in der vagen Hoffnung, diese später einmal monetarisieren zu können, widerspricht dem Geist des Gesetzes.

Zweckbindung

Daten, die beispielsweise zur Abwicklung eines Kaufvertrags erhoben wurden (Lieferadresse), dürfen nicht ohne Weiteres für andere Zwecke (z. B. Verkauf an Adresshändler oder aggressive Profilbildung für Fremdmaking) genutzt werden, es sei denn, es liegt eine gesonderte Einwilligung vor.

Neu und von technischer Tragweite sind die Prinzipien «Privacy by Design» (Datenschutz durch Technikgestaltung) und «Privacy by Default» (Datenschutz durch datenschutzfreundliche Voreinstellungen).

- **Privacy by Design** fordert, dass Datenschutzanforderungen bereits in der Entwicklungsphase von Produkten, Dienstleistungen und Systemen berücksichtigt werden. Wenn ein Finanzdienstleister eine neue App für das Spesenmanagement einführt, muss die Sicherheit der Daten Architekturprinzip sein, nicht ein nachträglich aufgesetztes Feature.
- **Privacy by Default** verlangt, dass die Voreinstellungen so gewählt sind, dass die Privatsphäre der Nutzer maximal geschützt ist. Der Nutzer soll nicht erst in tiefen Menüstrukturen nach Opt-out-Möglichkeiten suchen müssen. Für Softwareanbieter und Unternehmen bedeutet dies, dass Checkboxen für Newsletter oder Datenweitergaben standardmässig nicht angekreuzt sein dürfen.

6 Datensicherheit und Datenschutz

Vergleich: Schweizer DSG vs. EU-DSGVO

Obwohl eine Harmonisierung angestrebt wurde, existieren signifikante Unterschiede zwischen dem Schweizer und dem europäischen Recht. Diese Unterschiede sind für international tätige Schweizer Unternehmen von grosser Relevanz, da sie oft beide

Regelwerke parallel beachten müssen. Ein differenziertes Verständnis der Abweichungen verhindert teure Missverständnisse.

Abbildung 47 illustriert die wesentlichen strukturellen und inhaltlichen Divergenzen

Merkmal	EU-DSGVO	Schweizer DSG	Strategische Implikation für Unternehmen
Sanktionsadressat	Das Unternehmen (juristische Person).	Primär die natürliche Person (Verantwortlicher).	In der Schweiz zielt der Staat auf den Entscheidungsträger. Das Risiko ist persönlich und privatschädlich.
Bussenhöhe	Administrativbussen bis zu 4 % des weltweiten Jahresumsatzes oder 20 Mio. EUR.	Strafrechtliche Bussen bis zu CHF 250'000.	Während die EU-Bussen existenzbedrohend für die Firma sein können, sind CH-Bussen existenzbedrohend für den Manager privat.
Rechtsdogmatik	Verbot mit Erlaubnisvorbehalt (Alles verboten, ausser es ist erlaubt).	Erlaubnisprinzip mit Missbrauchsvorbehalt (Alles erlaubt, solange keine Persönlichkeitsverletzung).	Das Schweizer Recht gewährt mehr unternehmerische Freiheit, fordert aber höhere Eigenverantwortung in der Risikoabwägung.
Datenschutzbeauftragter	Pflicht für viele Unternehmen, streng reglementiert.	Grundsätzlich freiwillig, Pflicht nur in seltenen Hochrisikofällen.	Schweizer KMU haben organisatorischen Spielraum, sollten die Rolle aber zur Risikominimierung besetzen.
Meldepflicht (Breach)	Strikt innert 72 Stunden an die Behörde.	«So rasch als möglich» (keine starre Stundenfrist).	Die Schweiz gewährt mehr Flexibilität für eine gründliche Analyse vor der Meldung, verlangt aber dennoch unverzügliches Handeln.

Abbildung 47: Gegenüberstellung Schweizer DSG versus EU-DSGVO (eigene Darstellung)

6 Datensicherheit und Datenschutz

Abbildung 47 verdeutlicht, dass Schweizer Unternehmen sich nicht in falscher Sicherheit wiegen dürfen, nur weil das «Drohpotenzial» von Millionenbussen fehlt. Die persönliche Kriminalisierung von Fehlverhalten im Schweizer Recht schafft eine psychologische und reale Bedrohungslage für Führungskräfte, die in der EU in dieser Form nicht existiert. In der EU zahlt der Aktionär, in der Schweiz zahlt der Manager.

Organisatorische Massnahmen: Das operative Rückgrat der Compliance

Das Gesetz fordert nicht nur ein Ergebnis (sichere Daten), sondern schreibt auch den Weg dorthin vor. Unternehmen müssen «angemessene technische und organisatorische Massnahmen» (TOMs) implementieren. Während technische Massnahmen (Firewalls, Verschlüsselung, Pseudonymisierung) oft an die IT-Abteilung delegiert werden, sind organisatorische Massnahmen eine reine Führungsaufgabe. Sie definieren die Governance-Struktur, die Prozesse und die Kultur des Umgangs mit Daten. Ohne robuste organisatorische Verankerung nützt die beste Technologie nichts.

Governance und Verantwortlichkeiten

Datenschutz ist eine Querschnittsaufgabe, die nicht «nebenbei» erledigt werden kann. Das DSG empfiehlt explizit die Benennung eines betrieblichen **Datenschutzberaters** (Data Protection Advisor). Auch wenn dies für die meisten KMU keine gesetzliche Pflicht ist, stellt es eine dringende Handlungsempfehlung dar («Best Practice»).

Der Datenschutzberater fungiert als interne Anlaufstelle, überwacht die Einhaltung der Vorschriften, schult Mitarbeiter und berät die Geschäftsleitung bei neuen Projekten (z.B. Einführung eines neuen CRM-Systems). Um Interessenkonflikte zu vermeiden, ist die organisatorische Positionierung entscheidend. Der Datenschutzberater sollte nicht gleichzeitig IT-Leiter (der kontrolliert werden muss) oder CEO (der operative Entscheidungen trifft) sein. Eine direkte Berichtslinie an den Verwaltungs-

rat oder die Geschäftsleitung sichert die notwendige Unabhängigkeit und Eskalationsfähigkeit.

Unternehmen können die Kontaktdaten des Datenschutzberaters dem Eidgenössischen Datenschutz- und Öffentlichkeitsbeauftragten (EDÖB) melden. Dies bietet einen taktischen Vorteil: Bei einer Untersuchung oder Beschwerde kontaktiert die Behörde zuerst den Berater. Dies ermöglicht eine professionelle Kanalisierung der Kommunikation und verhindert oft, dass unbedachte Äusserungen von Mitarbeitern eine Situation eskalieren lassen.

Das Verzeichnis der Bearbeitungstätigkeiten (VBT)

Das VBT ist das zentrale Inventar aller Datenflüsse im Unternehmen. Es ist das Äquivalent zum Inventar in der Buchhaltung, nur dass statt Warenbeständen Datenbestände erfasst werden. Die Erstellung eines VBT ist oft der erste Moment der Wahrheit für ein Unternehmen, da hier «Schatten-IT» und unstrukturierte Prozesse ans Licht kommen.

Das Gesetz sieht zwar eine Ausnahme für KMU mit weniger als 250 Mitarbeitenden vor, diese greift jedoch nur, wenn das Unternehmen keine Daten bearbeitet, die ein hohes Risiko für die Persönlichkeit der Betroffenen bergen. Da fast jedes Unternehmen Lohnbuchhaltung (Gesundheitsdaten bei Krankheit, Bonitätsdaten, Religionszugehörigkeit für Quellensteuer usw.) oder Kundenprofile führt, ist die Ausnahme in der Praxis für die meisten Betriebe hinfällig.

Für jede Datensammlung (z.B. Lohnbuchhaltung, Zeiterfassung, Videoüberwachung, Newsletter-Datenbank) müssen folgende Parameter dokumentiert werden:

1. **Verantwortlichkeit:** Wer ist der «Owner» des Prozesses?
2. **Zweck:** Warum werden die Daten bearbeitet?
3. **Kategorien:** Welche Datenpunkte werden erhoben (Name, Adresse, IP, Bonität)?
4. **Empfänger:** An wen werden die Daten weitergegeben (Treuhänder, Cloud-Provider, Muttergesellschaft)?

6 Datensicherheit und Datenschutz

5. **Speicherdauer:** Wann werden die Daten gelöscht?
6. **Datensicherheit:** Wie werden die Daten geschützt (Zugriffskonzepte, Backups)?

Gerade für die Finanzabteilung ist dieses Verzeichnis essenziell, da hier oft unbemerkt lokale Excel-Listen mit Bonusberechnungen, exportierte Debitorenlisten oder manuelle Spesenabrechnungen existieren, die in keinem zentralen Sicherheitskonzept abgebildet sind. Das VBT zwingt zur Formalisierung dieser Prozesse.

Datenschutz-Folgenabschätzung (DSFA)

Wenn eine geplante Datenbearbeitung ein hohes Risiko für die Persönlichkeit oder die Grundrechte der betroffenen Personen mit sich bringt, muss zwingend vor Beginn der Bearbeitung eine DSFA durchgeführt werden. Dies ist ein formalisierter Risikomanagement-Prozess.

Im Finanzkontext ist dies besonders relevant, z. B. bei:

- **automatisierten Bonitätsprüfungen:** wenn Algorithmen entscheiden, ob ein Kunde auf Rechnung kaufen darf (Profiling)
- **Überwachung:** Einsatz von Software zur Analyse des Mitarbeiterverhaltens zur Betrugsprävention (Fraud Detection)
- **neuen Technologien:** Einführung von Biometrie für den Gebäudezugang oder KI-basierte Auswertung von Finanztransaktionen mit Personenbezug (z. B. Profiling im Retail-Banking, AML-/Fraud-Monitoring auf Kundenebene)

In der DSFA werden die Risiken systematisch analysiert, bewertet und Massnahmen definiert, um diese auf ein akzeptables Niveau zu senken. Dokumentiert das Unternehmen, dass trotz Massnahmen ein hohes Restrisiko verbleibt, muss der Eidgenössische Datenschutz- und Öffentlichkeitsbeauftragte (EDÖB) konsultiert werden – ein Szenario, das Unternehmen in der Regel vermeiden wollen.

Auftragsbearbeitung und Subunternehmer-Management

Moderne Unternehmen sind arbeitsteilig organisiert. Finanzabteilungen lagern Lohnbuchhaltungen an Treuhänder aus, nutzen Cloud-Speicher für Archive oder beauftragen Inkassobüros. Nach DSGVO bleibt das auftraggebende Unternehmen vollumfänglich für die Daten verantwortlich, auch wenn ein Dritter sie physisch bearbeitet. Man kann die Arbeit delegieren, aber nicht die Verantwortung.

Daraus resultieren organisatorische Pflichten:

- **Vertragliche Absicherung:** Es müssen Verträge zur Auftragsdatenbearbeitung (ADV-Verträge) abgeschlossen werden, die dem Dienstleister dieselben Pflichten auferlegen, die das Unternehmen selbst hat.
- **Kontrollpflicht:** Das Unternehmen muss sich vergewissern, dass der Dienstleister die Datensicherheit tatsächlich garantiert. Dies erfordert regelmässige Audits, Fragebögen oder das Einfordern von Zertifikaten (z. B. ISO 27001).
- **Auslandstransfer:** Besondere Vorsicht ist bei Auslagerung ins Ausland geboten. Es muss geprüft werden, ob das Zielland in Anhang 1 der Datenschutzverordnung (DSV) aufgeführt ist (angemessener Schutz). Die USA sind seit September 2024 nur für Unternehmen mit Swiss-US Data Privacy Framework-Zertifizierung gelistet. Für nicht zertifizierte US-Anbieter sowie Länder wie Indien oder China sind zusätzliche Garantien notwendig, wie der Abschluss von Standardvertragsklauseln (SCCs)³⁵ und die Durchführung eines Transfer Impact Assessments (TIA), um zu prüfen, ob lokale Gesetze (z. B. US CLOUD Act) den Zugriff auf Daten durch ausländische Behörden ermöglichen.

Meldewesen bei Verletzungen (Data Breach Notification)

Ein zentraler neuer Prozess ist die Meldepflicht bei Datensicherheitsverletzungen; anders als in der Vergangenheit, als Vorfälle häufig intern verblieben, verlangt das DSGVO ein hohes Mass an Transparenz.

³⁵Schweizerischer Bundesrat, 2022 (Datenschutzverordnung DSV, SR 235.11, Anhang 1); *economiesuisse*, o. J.; Stoiber, 2023.

6 Datensicherheit und Datenschutz

- **Meldung an den EDÖB:** Verletzungen der Datensicherheit, die voraussichtlich zu einem hohen Risiko für die betroffenen Personen führen, müssen «so rasch als möglich» gemeldet werden.
- **Information der Betroffenen:** Wenn es zu deren Schutz erforderlich ist (z.B. Passwort-Leaks, gestohlene Kreditkartendaten, Gefahr von Identitätsdiebstahl), müssen auch die Personen selbst informiert werden.

Organisatorische Herausforderung

Es muss ein interner «Emergency Response Plan» existieren. Ein Mitarbeiter in der Buchhaltung, der versehentlich eine Lohnliste an den falschen externen Empfänger mailt, muss wissen, wen er intern sofort informieren muss. Eine Kultur der Angst führt hier zu Vertuschung und somit zu Rechtsverstössen. Es braucht eine «Speak-Up»-Kultur, in der Fehler gemeldet werden können, um den Schaden zu begrenzen.

Datenklassifikation

Nicht alle Daten sind gleich schützenswert. Ein effektives Risikomanagement beginnt mit einer klaren Klassifikation, die dem Schutzbedarf der Daten entspricht. Im Finanz- und Rechnungswesen muss diese Klassifikation zwei Dimensionen abdecken: den Schutz der Persönlichkeit (DSG-Fokus) und den Schutz des Geschäftsbetriebs (Finanz- und Betriebsrisiko-Fokus).

- **Öffentliche Daten (geringes Risiko):** Daten, die bereits allgemein zugänglich sind
- **Interne allgemeine Daten (mittleres Risiko):** Nicht-öffentliche Kundendaten oder Lieferantendaten ohne sensitive Details
- **Besonders schützenswerte Daten gem. Art. 5 lit. c DSG (hohes Risiko):**³⁶
 - Daten über religiöse, weltanschauliche, politische oder gewerkschaftliche Ansichten oder Tätigkeiten
 - Daten über die Gesundheit, die Intimsphäre oder die Zugehörigkeit zu einer Rasse oder Ethnie

- genetische Daten, biometrische Daten, die eine natürliche Person eindeutig identifizieren
- Daten über verwaltungs- und strafrechtliche Verfolgungen oder Sanktionen
- Daten über Massnahmen der sozialen Hilfe

Das revidierte DSG erweitert diese Kategorie um Daten über administrative oder strafrechtliche Verfolgungen bzw. Sanktionen und Massnahmen der sozialen Hilfe. Im Finanzwesen fallen darunter Gesundheitsdaten von Mitarbeitenden, Daten zu Sozialhilfe oder Strafverfolgungen sowie Lohnabrechnungen, soweit diese indirekt Gesundheits-, Gewerkschafts- oder Sozialhilfeeintragungen offenlegen (z.B. IV-Taggelder, Krankentaggelder, Gewerkschaftsabzüge). Die Bearbeitung solcher Daten löst erhöhte Schutzpflichten aus.

- **Geschäftskritische Daten (operationelles Risiko):** Daten, deren Verlust oder Manipulation die operationelle Resilienz unmittelbar gefährdet, wie ERP-Stammdaten, Zugänge zu Zahlungssystemen oder proprietäre Handelsinformationen

Nur mit der richtigen Klassifizierung kann das Management die angemessenen technischen und organisatorischen Massnahmen (TOMs) implementieren und die Einhaltung der gesetzlichen Anforderungen gewährleisten.

6.3 Datensicherheit

Die Datensicherheit befasst sich mit der technischen und organisatorischen Umsetzung des Schutzes für den gesamten Datenbestand eines Unternehmens. Während der Datenschutz die rechtliche Zulässigkeit regelt, stellt die Datensicherheit die operative Integrität, Vertraulichkeit und Verfügbarkeit der Informationen sicher. Im Finanzwesen ist sie essenziell, um die Vorgaben der Geschäftsbücherverordnung (GeBüV) zu erfüllen und die Verlässlichkeit der Rechnungslegung zu garantieren.

Unternehmen sind gesetzlich verpflichtet, angemessene technische und organisatorische Massnahmen (TOMs) zu implementieren, die dem Risikoprofil der

³⁶Bundesversammlung der Schweizerischen Eidgenossenschaft, 2020a (Bundesgesetz über den Datenschutz DSG, SR 235.1, Art. 5 lit. c).

6 Datensicherheit und Datenschutz

verarbeiteten Daten entsprechen. Dazu gehören Verschlüsselungstechnologien, Zugriffskontrollen und Protokollierungssysteme, die sicherstellen, dass Finanzdaten vor unbefugter Veränderung, Verlust oder Diebstahl geschützt sind, unabhängig davon, ob die Bedrohung von aussen oder innen stammt.

Definition und Zielsetzung in der Finanzarchitektur

Die Datensicherheit (Data Security) befasst sich mit dem Schutz von Datenbeständen vor unbefugtem Zugriff, Manipulation und Verlust durch die Implementierung technischer und organisatorischer Massnahmen. Während der Datenschutz die rechtliche Zulässigkeit der Datenbearbeitung regelt, gewährleistet die Datensicherheit die technische Durchsetzbarkeit dieser Regeln.

Für die Finanzbuchhaltung und das Controlling ist die Datensicherheit eine fundamentale Voraussetzung für die Ordnungsmässigkeit der Buchführung. Die Zielsetzung orientiert sich an der **CIA-Triade** (Vertraulichkeit, Integrität, Verfügbarkeit), welche im Schweizer Kontext durch die Anforderungen der Geschäftsbücherverordnung (GeBüV) sowie die Vorgaben der FINMA an die Integrität von Daten konkretisiert wird:

Vertraulichkeit (Confidentiality):

- *Definition:* Informationen sind ausschliesslich autorisierten Personen oder Systemen zugänglich.
- *Finanzspezifische Anforderung:* Zugriffsbeschränkungen müssen granulare Berechtigungskonzepte abbilden, die konsequent dem Need-to-know-Prinzip folgen: Mitarbeitende erhalten ausschliesslich Zugriff auf Finanzdaten, welche für die Erfüllung der konkreten Funktion zwingend erforderlich sind.

Integrität (Integrity):

- *Definition:* Daten müssen korrekt, vollständig und vor unautorisierter Veränderung geschützt sein.
- *Finanzspezifische Anforderung:* Gemäss Art. 3 GeBüV³⁷ muss sichergestellt sein, dass Buchun-

gen nach ihrer Festschreibung nicht mehr verändert werden können, ohne dass dies protokolliert wird. Die Datenintegrität ist die Basis für verlässliche Finanzberichte und Audits.

Verfügbarkeit (Availability):

- *Definition:* Daten und die zugehörigen Systeme stehen zum benötigten Zeitpunkt und in der geforderten Performance zur Verfügung.
- *Finanzspezifische Anforderung:* Kritische Prozesse wie der Monatsabschluss, die Lohnlauf-Ausführung oder das Liquiditätsmanagement sind zeitkritisch. Ein Ausfall der Verfügbarkeit kann unmittelbare finanzielle Schäden (Verzugszinsen, Vertragsstrafen) nach sich ziehen.

Technische und Organisatorische Massnahmen (TOMs)

Das Datenschutzgesetz (DSG) sowie internationale Standards (ISO 27001) fordern die Umsetzung geeigneter technischer und organisatorischer Massnahmen (TOMs). Die Angemessenheit dieser Massnahmen bemisst sich nach dem Schutzbedarf der Daten und den möglichen Risiken für die betroffenen Personen und das Unternehmen.

Technische Massnahmen (Infrastruktur- und Datenebene)

Verschlüsselung und Kryptografie:

- *Data at Rest:* Datenbanken (z.B. SAP HANA, SQL) sowie mobile Endgeräte und Datenträger müssen verschlüsselt werden. Dies verhindert, dass bei physischem Diebstahl von Hardware auf die Daten zugegriffen werden kann.
- *Data in Transit:* Jegliche Übermittlung von Finanzdaten (z.B. Zahlungsfiles im pain.001-Format an Banken, Steuerdaten via ELM) muss über gesicherte Kanäle erfolgen (z.B. TLS 1.3, EBICS-Protokolle).

Data Loss Prevention (DLP):

- DLP-Systeme überwachen Datenströme auf Muster (z.B. Kreditkartennummern, IBANs, AHV-Nummern) und verhindern den unautorisierten Abfluss.
- *Use Case:* Ein DLP-System blockiert den Versuch

³⁷ Schweizerischer Bundesrat, 2002 (Verordnung über die Führung und Aufbewahrung der Geschäftsbücher GeBüV, SR 221.431, Art. 3); aR solutions, o. J.; swiDOC, o. J.

6 Datensicherheit und Datenschutz

eines Mitarbeitenden, eine Excel-Liste mit dem gesamten Kundenstamm auf einen privaten USB-Stick zu kopieren oder per Webmail zu versenden.

Protokollierung (Logging & Monitoring):

- Zur Erfüllung der GeBüV (Art. 8) und der Nachvollziehbarkeit müssen Zugriffe (Lesen) und Mutationen (Schreiben/Löschen) protokolliert werden.
- *Anforderung:* Log-Dateien müssen manipulationssicher aufbewahrt werden (z.B. auf WORM-Speichern), um sicherzustellen, dass privilegierte Administratoren ihre Spuren nach einer unautorisierten Handlung nicht löschen können.

Organisatorische Massnahmen (Prozesse und Governance)

Technische Lösungen sind wirkungslos ohne flankierende organisatorische Regelungen. Im Finanzbereich stehen hier Kontrollmechanismen im Vordergrund.

Funktionstrennung (Segregation of Duties - SoD):

- SoD verhindert, dass eine einzelne Person einen kritischen Prozess vollständig selbst durchführen kann, was das Risiko von Betrug und Fehlern minimiert.
- *SoD-Matrix:* Unternehmen sollten eine Matrix definieren, die inkompatible Rollen im ERP-System ausweist. Beispiel: Die Rolle «Anlage von Kreditorenstammdaten» darf nicht mit der Rolle «Zahlungsfreigabe» kombiniert werden.

Vier-Augen-Prinzip:

- Für Transaktionen ab einem definierten Risikowert (z.B. Zahlungen grösser CHF 5000) ist technisch die Freigabe durch eine zweite autorisierte Person zu erzwingen. Dies ist eine effektive Massnahme gegen CEO-Fraud und interne Veruntreuung.

Identitäts- und Zugriffsmanagement (IAM) & Offboarding:

- Der Austrittsprozess ist ein kritischer Sicherheitsfaktor. Es muss organisatorisch sichergestellt sein, dass Zugriffsrechte auf ERP-Systeme und Bankvollmachten exakt zum Zeitpunkt des Austritts (oder der Freistellung) entzogen werden. Veraltete Accounts («Geister-User») stellen ein signifikantes Sicherheitsrisiko dar.

Privileged Access Management (PAM)

In Finanzabteilungen und der unterstützenden IT existieren Benutzerkonten mit weitreichenden Rechten (z.B. ERP-Administratoren, Treasury-Leiter). Diese Konten sind primäre Angriffsziele.

- **Konzept:** Administratoren arbeiten im Tagesgeschäft mit einem Standard-Konto. Für kritische administrative Tätigkeiten nutzen sie temporär privilegierte Accounts, die über eine PAM-Lösung verwaltet werden.
- **Session Recording:** Bei hochkritischen Eingriffen (z.B. Änderung der Bankverbindung des eigenen Unternehmens) können Sitzungen aufgezeichnet werden, um eine lückenlose forensische Beweisführung zu ermöglichen.

6.4 Cyber Security

Cyber Security konzentriert sich auf den aktiven Schutz der IT-Infrastruktur, Netzwerke und Endgeräte vor böswilligen Angriffen aus dem Cyberraum. Die Bedrohungslage ist durch eine zunehmende Professionalisierung der Angreifer gekennzeichnet, die gezielt Schwachstellen in Systemen oder menschliche Faktoren ausnutzen, um Zugang zu Unternehmensnetzwerken zu erlangen. Für Finanzabteilungen stehen dabei Risiken wie Ransomware, CEO-Fraud und gezielte Angriffe auf den Zahlungsverkehr im Vordergrund.

Eine effektive Cyber-Security-Strategie beschränkt sich nicht auf technische Abwehrmechanismen wie Firewalls, sondern erfordert einen ganzheitlichen Ansatz, der auch die Erkennung von Angriffen und die Sensibilisierung der Mitarbeitenden umfasst. Angesichts der dynamischen Entwicklung von Angriffsmethoden, einschliesslich des Einsatzes von künstlicher Intelligenz durch Cyberkriminelle, ist ein kontinuierliches Monitoring und die Anpassung der Schutzmassnahmen unerlässlich.

6 Datensicherheit und Datenschutz

Die Bedrohungslage: Ein industrialisierter Markt

Die aktuelle Bedrohungslage in der Schweiz: Hackerangriffe sind das dominante Risiko. Mehr als acht von zehn Schweizer Unternehmen (81,6%) erwarten einen Anstieg von Cyberangriffen. Gleichzeitig verfügt mehr als ein Drittel (32,7%) der Unternehmen über keinerlei formelle Risikomanagementstrukturen. Diese Diskrepanz zwischen wahrgenommener Bedrohung und unzureichender Vorbereitung schafft ein kritisches Risiko³⁸. Hinzu kommt, dass neue generative KI-Technologien die Erstellung gefälschter Profile oder Zahlungsbelege vereinfachen, was die Gefahr von Betrugsversuchen drastisch erhöht.

Dabei hat sich die Natur der Angriffe gewandelt. Es geht längst nicht mehr nur um technische Lücken in der Firewall. Moderne Angriffe zielen zunehmend auf den Faktor Mensch: Durch psychologische Manipulation (Social Engineering) oder täuschend echte E-Mails (Phishing) versuchen Angreifer, die Mitarbeitenden im Rechnungswesen dazu zu bringen, Überweisungen zu tätigen oder Zugangsdaten preiszugeben.

Cyberkriminalität hat sich zu einem gewinnorientierten Geschäftsmodell entwickelt. Die Barrieren für Angreifer sind so niedrig wie nie zuvor.

- **Industrialisierung:** Über Modelle wie Ransomware-as-a-Service (RaaS) können Kriminelle fertige Angriffsinfrastrukturen mieten. Sie verfügen über «Kundendienst»-Abteilungen, die Opfern helfen, Bitcoins zu kaufen, um Lösegelder zu bezahlen.
- **Supply Chain Risk:** Angreifer attackieren zunehmend nicht das Unternehmen direkt, sondern dessen Software-Dienstleister (z. B. Anbieter von Buchhaltungs- oder Lohnsoftware). Ist der Dienstleister kompromittiert, haben die Angreifer einen «Generalschlüssel» zu den Daten deren Mandanten.

Die Angreifer: Wer bedroht uns?

Für den CFO ist es wichtig zu verstehen, wer auf der anderen Seite steht, um die Risiken richtig einzuschätzen:

- **Organisierte Kriminalität:** Der Grossteil der Angriffe ist finanziell motiviert. Diese Gruppen agieren wie Wirtschaftsunternehmen, optimieren ihren ROI und suchen gezielt nach solventen Opfern, deren Ausfallkosten (z. B. bei einem Produktionsstopp) so hoch sind, dass sie bereitwillig Lösegeld zahlen.
- **Staatliche Akteure (Nation State Actors):** Insbesondere für international tätige Schweizer Unternehmen relevant. Diese Gruppen betreiben Wirtschaftsspionage oder suchen Devisenbeschaffung (z. B. nordkoreanische Gruppen, die gezielt Finanzinstitute oder Kryptobörsen angreifen).
- **Insider (Insider Threat):** Oft unterschätzt, aber fatal: Frustrierte Mitarbeitende oder solche, deren Zugangsdaten unbemerkt gestohlen wurden, können Sicherheitsmechanismen von innen umgehen.

Spezifische Bedrohungsszenarien für Finance & Accounting

Die Finanzabteilung ist das «Kronjuwel» für Cyberkriminelle. Hier laufen die Zahlungsströme zusammen, hier liegt die Verfügungsgewalt über liquide Mittel. Die Angriffe des Jahres 2025 zielen dabei weniger auf technische Exploits der Buchhaltungssoftware, sondern neuerdings auch auf die Manipulation der menschlichen Entscheidungsträger und Prozesse (Business Email Compromise – BEC).

CEO-Fraud 2.0: Die Gefahr der synthetischen Identitäten

Der klassische CEO-Fraud, bei dem eine gefälschte E-Mail im Namen des Geschäftsführers eine dringende Auslandsüberweisung fordert, ist bekannt. 2025 hat sich diese Bedrohung durch «Deepfake»-Technologie qualitativ verändert.

³⁸ Hochschule Luzern HSLU / Schweizerischer Gewerbeverband sgV, 2025.

6 Datensicherheit und Datenschutz

Deepfake-Video-Calls und Voice-Cloning

In einem vom BACS dokumentierten Fall wurden Finanzverantwortliche nicht mehr nur per Mail kontaktiert, sondern zu Video-Konferenzen eingeladen. In diesen Calls sahen sie ihren Vorgesetzten und hörten seine Stimme – beides jedoch synthetisch generiert durch KI.³⁹

- **Technologie:** Angreifer nutzen öffentlich verfügbares Audio- und Videomaterial (z. B. von YouTube, Podcasts, Firmenwebseiten), um KI-Modelle zu trainieren. Wenige Minuten Quellmaterial reichen oft aus, um eine täuschend echte Stimme zu klonen («Voice Cloning»).
- **Psychologie:** Die visuelle und auditive Bestätigung durch den vermeintlichen Chef hebt die üblichen Sicherheitsmechanismen aus («Ich habe ihn ja gesehen!»). In Kombination mit einem konstruierten Dringlichkeitsszenario (z. B. eine geheime Firmenübernahme) wird der Mitarbeiter zur Umgehung von Kontrollprozessen gedrängt.

Präventionsstrategien für die Buchhaltung

Gegen diese technologisch hochgerüsteten Angriffe hilft keine technische Firewall, sondern nur eine «Human-Firewall» und rigide Prozesse.

- **Prozess-Integrität:** Zahlungen über einem bestimmten Schwellenwert dürfen niemals nur auf Basis elektronischer Kommunikation (E-Mail, Video, Telefon) ausgelöst werden.
- **Out-of-Band-Verifikation:** Ein Rückruf unter einer intern bekannten Nummer (nicht die Nummer, die im Display angezeigt wird oder in der Mail steht) oder ein persönliches Gespräch sind zwingend erforderlich.
- **Codewörter:** Für Notfälle können vorab vereinbarte Codewörter genutzt werden, die physisch hinterlegt sind und in keinem digitalen Kanal kommuniziert wurden.

Quishing: Die Lücke in der QR-Rechnung

Mit der flächendeckenden Einführung der QR-Rechnung in der Schweiz hat sich ein neuer Angriffsvektor etabliert: «Quishing» (QR-Code-Phishing).

Der Mechanismus des Betrugs: Kriminelle versenden Rechnungen (per E-Mail oder sogar per Post), die legitim aussehen und oft reale Geschäftsbeziehungen imitieren. Der QR-Code auf der Rechnung ist jedoch manipuliert.

- **Szenario A (Phishing):** Der Scan des Codes führt nicht zur Banking-App, sondern auf eine gefälschte Webseite, die Login-Daten abgreift.
- **Szenario B (Zahlungsumleitung):** Der QR-Code enthält die korrekten Rechnungsdaten (Betrag, Referenz), aber eine falsche IBAN oder einen falschen Begünstigten. Da viele Nutzer und ERP-Systeme primär auf den erfolgreichen Scan vertrauen und die Details auf dem Bildschirm nicht mehr akribisch prüfen, wird die Zahlung an ein Geldwäschekonto ausgelöst.

Risiken bei der automatisierten Verarbeitung

Besonders gefährdet sind Unternehmen mit hohem Rechnungsvolumen und automatisierter Belegverarbeitung. Wenn das ERP-System (z. B. Abacus, Bexio, SAP) den QR-Code ausliest, muss sichergestellt sein, dass eine Stammdatenvalidierung stattfindet.

- **Abweichungs-Check:** Stimmt die im QR-Code kodierte IBAN mit der im Kreditorenstamm hinterlegten IBAN überein? Wenn nein, muss das System die Zahlung zwingend blockieren und zur manuellen Prüfung vorlegen. Banken prüfen bei Kleinbeträgen oft nicht die Übereinstimmung von Namen und IBAN, was die Verantwortung vollständig auf den Zahler verlagert.

³⁹ Bundesamt für Cybersicherheit BACS, 2024.

6 Datensicherheit und Datenschutz

Rechnungsmanipulation (Invoice Redirection Fraud)

Bei dieser Betrugsform hacken Kriminelle nicht das Unternehmen selbst, sondern dessen Lieferanten. Sie überwachen den E-Mail-Verkehr und greifen ein, wenn eine legitime Rechnung versendet wird. Die Rechnung wird abgefangen, die Bankverbindung manipuliert und die E-Mail dann weitergeleitet.

- **Erkennungsproblem:** Da die Rechnung erwartet wird, der Betrag stimmt und der Absender (technisch gesehen) korrekt ist, ist dieser Betrug extrem schwer zu erkennen.
- **Statistik:** Das BACS verzeichnete 2025 eine Zunahme dieser Fälle⁴⁰, oft begünstigt durch unzureichende E-Mail-Sicherheit (fehlendes 2-Faktor-Login) bei Zulieferern.

Verteidigung: Das Three Lines of Defense Modell

Angesichts dieser industrialisierten Bedrohungslage reicht es nicht mehr aus, die Verantwortung allein an die IT-Abteilung zu delegieren. Ein effektives Abwehrdispositiv erfordert eine klare Governance-Struktur. Das 3-Linien-Modell (Three Lines of Defense) bietet hierfür den idealen Rahmen, um Verantwortlichkeiten im Accounting klar zu regeln.

Verteidigungslinie 1: Operatives Management (Daily Business)

Die Mitarbeitenden in der Buchhaltung und im Controlling sind die wichtigste Barriere gegen Angriffe.

- **Interne Kontrollsysteme (IKS):** Konsequente Anwendung des Vier-Augen-Prinzips bei allen Zahlungen über einer bestimmten Schwelle. Dies verhindert, dass ein einzelner kompromittierter Account oder ein getäuschter Mitarbeiter (CEO-Fraud) Gelder abfliessen lassen kann.
- **Identitätsschutz:** Nutzung von Multi-Faktor-Authentifizierung (MFA) und Passwort-Managern als nicht verhandelbarer Standard für den Zugriff auf ERP- und Bank-Systeme.

- **Awareness:** Kritische Prüfung von E-Mails (Absender, Links) und Verifikation von geänderten Kontodaten durch Rückruf beim Lieferanten (über eine bekannte Nummer, nicht die in der Mail angegebene).

Verteidigungslinie 2: Risikomanagement & Compliance (Überwachung)

Diese Ebene erstellt die Regeln und überwacht deren Einhaltung. In KMU wird diese Rolle oft vom CFO oder einem externen Beauftragten wahrgenommen.

- **Richtlinienkompetenz:** Definition von klaren Weisungen (Policies), z. B. wie mit mobilen Geräten gearbeitet werden darf oder welche Cloud-Dienste (Schatten-IT) verboten sind.
- **Vendor Risk Management:** Überprüfung der Sicherheitszertifikate (z. B. ISO 27001) von Software-Lieferanten und Treuhändern vor Vertragsabschluss.
- **Notfallplanung (BCM):** Regelmässiges Testen der Backups und Aktualisierung der Notfallpläne («Was tun wir, wenn das ERP morgen verschlüsselt ist?»).

Verteidigungslinie 3: Interne Revision / Audit (Unabhängige Prüfung)

Die dritte Linie prüft unabhängig, ob die ersten beiden Linien funktionieren.

- **Penetration-Tests:** Beauftragung von ethischen Hackern, die versuchen, kontrolliert in das System einzudringen, um Lücken zu finden, bevor es Kriminelle tun.
- **Auditierung:** Prüfung, ob Zugriffsrechte noch aktuell sind (z. B. ob ausgetretene Mitarbeiter im System gesperrt wurden) und ob die definierten IKS-Prozesse tatsächlich gelebt werden.

⁴⁰Bundesamt für Cybersicherheit BACS, 2024.

6 Datensicherheit und Datenschutz

Cyber Security Awareness

Cyber Security Awareness ist eine kontinuierliche organisatorische Massnahme. Die FINMA hat festgestellt, dass bei vielen beaufsichtigten Instituten das **Cyber-Training und -Bewusstsein** verbesserungswürdig sind. Nur regelmässig geschulte Mitarbeitende können das Risiko mindern, das durch das Vertrauen auf rein technische Massnahmen entsteht.

Die Schulungsinhalte müssen auf die spezifischen Tätigkeiten im Finanz- und Rechnungswesen zugeschnitten sein:

- **Umgang mit sensiblen Daten:** Mitarbeitende müssen wissen, wie besonders schützenswerte Personendaten (z.B. Gesundheitsdaten, Transaktionsdetails) behandelt, pseudonymisiert oder verschlüsselt werden müssen.
- **Erkennung von Manipulationen im Zahlungsverkehr:** Ein zentrales Thema ist die Verifizierung von Zahlungsanweisungen und die Erkennung von Manipulationen (z.B. bei Änderungen von Bankverbindungen).
- **Betroffenenrechte-Prozesse:** Mitarbeitende müssen geschult sein, wie sie Auskunftsanfragen über Daten von Privatpersonen (DSG) schnell und unter Einhaltung der Identitätsprüfungsvorschriften behandeln.
- **Verhaltensregeln im Notfall:** Jede Person im Unternehmen muss wissen, welche internen Kontakte bei Anzeichen eines Cyberangriffs sofort zu informieren sind. Klare Meldekette sind die Voraussetzung für eine «so rasch wie möglich» Meldung an den EDÖB (Art. 24 DSG) und die frühzeitige Warnung an die FINMA (Art. 29 Abs. 2 FINMAG).

6.5 Incident Management

Incident Management bezeichnet den strukturierten Prozess zur Reaktion auf Sicherheitsvorfälle, die die Vertraulichkeit, Integrität oder Verfügbarkeit von Daten und Systemen gefährden. Ziel ist es, im Ernstfall schnell und koordiniert zu handeln, um den Schaden für das Unternehmen zu begrenzen und die Wiederherstellungszeit zu minimieren. Dies erfordert vordefinierte Abläufe, klare Verantwortlichkeiten und eine vorbereitete Krisenkommunikation.

Besondere Bedeutung kommt dem Incident Management durch die strengen regulatorischen Meldepflichten in der Schweiz zu. Finanzinstitute und Unternehmen müssen in der Lage sein, Vorfälle zeitnah zu identifizieren, zu klassifizieren und bei Bedarf fristgerecht an Behörden wie die FINMA oder den EDÖB zu melden. Eine verzögerte oder unkoordinierte Reaktion kann rechtliche Sanktionen und Reputationsschäden erheblich verschärfen.

Indikatoren für Finanz-Incidents:

- Eingang von Rechnungen mit manipulierten Zahlungsverbindungen (Payment Diversion Fraud)
- Systemmeldungen über unautorisierte Zugriffsversuche auf Treasury-Systeme
- Verschlüsselung von Buchhaltungsdaten durch Ransomware
- Meldungen von Partnern oder Kunden über ungewöhnliche Kommunikation (Verdacht auf kompromittierte E-Mail-Konten)

Eine klare Klassifizierungsmatrix (z.B. Schweregrad 1 bis 4) hilft, die richtigen Eskalationspfade (z.B. Einberufung des Krisenstabs) zeitgerecht zu aktivieren.

6 Datensicherheit und Datenschutz

Der Incident-Response-Prozess (Phasenmodell)

Ein professionelles Incident Management folgt einem standardisierten Zyklus (in Anlehnung an NIST SP 800-61⁴¹ oder ISO/IEC 27035⁴²):

1. Vorbereitung (Preparation):

- Etablierung eines Incident-Response-Teams (IRT) mit Vertretern aus IT, Legal, Kommunikation und Finance/CFO
- Bereitstellung von «Out-of-Band»-Kommunikationskanälen (z. B. sichere Messenger, dedizierte Mobiltelefone), um die Kommunikation aufrechtzuerhalten, falls das interne E-Mail-System kompromittiert ist

2. Erkennung und Analyse (Detection & Analysis):

- Identifikation des Vorfalls und Bestimmung des Ausmasses (Scope). Welche Daten (CID, Mitarbeiterdaten) und Systeme sind betroffen?
- Forensic Readiness: In dieser Phase dürfen betroffene Systeme nicht voreilig neu gestartet werden, da dies flüchtige Spuren im Arbeitsspeicher (RAM) löscht, die für die Beweissicherung und Versicherungsansprüche essenziell sind.

3. Eindämmung (Containment):

- Isolierung betroffener Systeme vom Netzwerk, um die Ausbreitung (Lateral Movement) von Schadsoftware zu unterbinden; Sperrung kompromittierter Benutzerkonten und temporäre Deaktivierung von Schnittstellen zu Banken

4. Beseitigung (Eradication):

- Entfernung der Ursache (Malware, Rootkits), Schliessen der ausgenutzten Sicherheitslücken und Bereinigung der Systeme

5. Wiederherstellung (Recovery):

- Wiederinbetriebnahme der Systeme aus verifizierten Backups; Validierung der Datenintegrität durch die Fachbereiche (z. B. Abstimmung der Hauptbuchkonten), bevor der operative Betrieb wieder freigegeben wird

6. Nachbereitung (Lessons Learned):

- Analyse des Vorfalls zur Optimierung der Sicherheitsarchitektur und Prozesse

Regulatorisches Meldewesen und Fristen

In der Schweiz bestehen parallele Meldepflichten, deren Einhaltung für die Compliance und Haftungsminimierung zentral ist.

Meldung an die FINMA (Art. 29 Abs. 2 FINMAG)

- **Adressaten:** beaufsichtigte Institute (Banken, Versicherungen, Vermögensverwalter)
- **Auslöser:** Vorfälle von «wesentlicher Bedeutung». Dies umfasst Cyber-Attacken, die kritische Funktionen beeinträchtigen oder die Schutzziele kritischer Daten verletzen.
- **Prozess:**
 - Sofortmeldung: innert 24 Stunden nach Entdeckung an den zuständigen Key Account Manager der FINMA
 - Detaillierte Meldung: innert 72 Stunden über die Erhebungs- und Gesuchsplattform (EHP) inklusive einer Analyse der Ursachen (Root Cause Analysis⁴³)

⁴¹National Institute of Standards and Technology NIST, 2025.

⁴²International Organization for Standardization, 2023 (ISO/IEC 27035-1:2023).

⁴³Eidgenössische Finanzmarktaufsicht FINMA, 2024 (Aufsichtsmittteilung 03/2024).

6 Datensicherheit und Datenschutz

Meldung an den EDÖB (Art. 24 DSG)

- **Adressaten:** alle datenbearbeitenden Unternehmen (Verantwortliche)
- **Auslöser:** Verletzungen der Datensicherheit, die voraussichtlich zu einem **hohen Risiko** für die Persönlichkeit oder die Grundrechte der betroffenen Personen führen (z. B. Abfluss von Gesundheitsdaten, biometrischen Daten, einer grossen Anzahl Kundendaten oder Lohnabrechnungen mit Hinweisen auf Sozialhilfe oder Krankheit)
- **Frist:** «so rasch als möglich». In der Praxis bedeutet dies unverzüglich nach Erlangung eines ausreichenden Kenntnisstandes zur Risikoeinschätzung (i. d. R. innert 72 Stunden).
- **Prozess:** Meldung via «DataBreach»-Portal des Bundes. Falls zum Schutz der Personen erforderlich, müssen diese ebenfalls informiert werden.

Meldung an das BACS

- **Adressaten:** Betreiber kritischer Infrastrukturen (seit April 2025 obligatorisch, für andere empfohlen⁴⁴)
- **Frist:** innert 24 Stunden nach Entdeckung des Cyberangriffs

6.6 Business Continuity Management (BCM)

Business Continuity Management (BCM) zielt auf die Aufrechterhaltung kritischer Geschäftsprozesse während und nach einer schwerwiegenden Störung ab. Der Fokus hat sich dabei von der reinen IT-Wiederherstellung hin zur operationellen Resilienz verschoben. Es geht primär darum, die Erbringung essenzieller Dienstleistungen für Kunden und Marktpartner auch unter widrigen Umständen sicherzustellen, selbst wenn reguläre Systeme temporär nicht verfügbar sind.

Basis eines robusten BCM ist die Identifikation kritischer Funktionen und die Festlegung von Toleranzgrenzen für deren Ausfall. Durch eine Kombination aus technischer Redundanz, wie Georedundanz (zur

Erläuterung siehe unten) und unveränderbaren Backups, sowie organisatorischen Notfallplänen wird die Überlebensfähigkeit des Unternehmens gesichert. Regelmässige Tests und Simulationen stellen sicher, dass diese Pläne im Krisenfall nicht nur auf dem Papier existieren, sondern auch operativ greifen.

Business Impact Analyse (BIA) und Kritikalitätsbewertung

Die Basis des BCM bildet die Business-Impact-Analyse (BIA). Hierbei werden nicht primär IT-Systeme, sondern Geschäftsprozesse hinsichtlich ihrer Zeitkritikalität bewertet.

- **Identifikation kritischer Funktionen:** Prozesse, deren Ausfall die Stabilität des Instituts bedroht oder wesentliche Schäden für Kunden verursacht
 - *Beispiele:* Zahlungsverkehr (SIC/SEPA), Liquiditätsbewirtschaftung, Börsenhandel
- **Impact Tolerance (Toleranzgrenze):** Für jede kritische Funktion muss definiert werden, wie lange eine Unterbrechung maximal andauern darf, bevor irreversible Schäden eintreten (z. B. Reputationsverlust, regulatorische Sanktionen). Diese Toleranzgrenze diktiert die technischen Parameter:
 - *RTO (Recovery Time Objective):* die maximal zulässige Zeit bis zur Wiederherstellung des Prozesses
 - *RPO (Recovery Point Objective):* der maximal tolerierbare Datenverlust (z. B. «0 Transaktionen» im Zahlungsverkehr).

Kontinuitätsstrategien und Workarounds

Um die definierten Toleranzgrenzen einzuhalten, sind redundante Systeme und alternative Prozesse notwendig.

⁴⁴Bundesversammlung der Schweizerischen Eidgenossenschaft, 2020b (Informationssicherheitsgesetz ISG, SR 128, Art. 74a).

6 Datensicherheit und Datenschutz

Technische Redundanz und Datensicherung

- **Georedundanz:** Spiegelung kritischer Systeme in ein physisch getrenntes Rechenzentrum (Hot-Standby), um bei lokalen Ereignissen handlungsfähig zu bleiben
- **Immutable Backups:** Angesichts der Bedrohung durch Ransomware, die auch Backups verschlüsselt, ist der Einsatz von unveränderbaren Speichern (WORM-Technologie) oder isolierten «Cyber Vaults» (Air-Gapped) essenziell. Diese dienen als letzte Verteidigungslinie zur Wiederherstellung eines sauberen Datenbestandes.

Organisatorische Notfallprozesse (Manual Workarounds)

Ein robustes BCM beinhaltet detaillierte Pläne für den Betrieb ohne reguläre IT-Infrastruktur.

- **Zahlungsverkehr:** Existieren Prozesse und Berechtigungen, um Grosszahlungen via Notfall-Fax oder dedizierte Terminals direkt bei der Nationalbank oder Korrespondenzbanken auszulösen? Lohnbuchhaltung: Stehen «Clean Laptops» (Geräte, die nie mit dem kompromittierten Netzwerk verbunden waren) und aktuelle Datensätze bereit, um im Notfall Lohnzahlungen manuell via E-Banking zu erfassen?
- **Liquiditätsstatus:** Wie wird der Cash-Status ermittelt, wenn das Treasury-Management-System nicht verfügbar ist? (z.B. telefonische Saldenabfrage).

Validierung durch Tests und Übungen

Ein BCM-Konzept ist nur so effektiv wie seine Validierung. Die FINMA fordert regelmässige Tests basierend auf «schwerwiegenden, aber plausiblen» Szenarien.

- **Tabletop Exercises:** Simulationen auf Management-Ebene zur Überprüfung von Entscheidungswegen, Kommunikationsketten und Kompetenzregelungen in Krisensituationen
- **Disaster Recovery Tests:** Technisches Umschalten auf Ausweichsysteme (Failover) zur Verifizierung der RTO-Vorgaben

- **Restore Tests:** Regelmässige Prüfung, ob Backups nicht nur technisch lesbar, sondern auch inhaltlich konsistent wiederhergestellt werden können

6.7 Der CFO-Blickwinkel

Vom Kostenfaktor zum strategischen Risikomanagement

Die Analyse der Bereiche Datenschutz, Datensicherheit, Cyber Security und Resilienz verdeutlicht einen fundamentalen Wandel: Informationssicherheit ist keine reine IT-Disziplin mehr, sondern ein zentraler Bestandteil des unternehmerischen Risikomanagements. Für den CFO bedeutet dies, dass Investitionen in Sicherheitsarchitekturen nicht als blosser Kostenfaktor («Cost of Doing Business»), sondern als Massnahme zum Bilanzschutz und zur Wertsicherung («Asset Protection») bewertet werden müssen.

In einer digitalisierten Finanzwelt korreliert die Stabilität der IT-Systeme direkt mit der Liquidität und Handlungsfähigkeit des Unternehmens. Ein erfolgreicher Ransomware-Angriff oder ein gross angelegter Betrugsfall (CEO-Fraud) bedroht nicht nur die kurzfristige Profitabilität, sondern kann durch Reputationsverlust und Kundenabwanderung den Unternehmenswert nachhaltig schädigen. Der CFO muss daher sicherstellen, dass Cyber-Risiken quantifiziert und im internen Kontrollsystem (IKS) adäquat abgebildet sind.

Die neue Dimension der Haftung

Mit dem revidierten Datenschutzgesetz (DSG) und den verschärften Vorgaben der FINMA hat sich das Haftungsrisiko massiv auf die Führungsebene verlagert. Die Möglichkeit persönlicher Bussen von bis zu CHF 250 000 für vorsätzliche Pflichtverletzungen von Entscheidungsträgern schafft eine neue Realität: Compliance ist Chefsache.

6 Datensicherheit und Datenschutz

Der CFO trägt hierbei eine doppelte Verantwortung:

1. **Ressourcenallokation:** Er muss sicherstellen, dass ausreichende Mittel für angemessene technische und organisatorische Massnahmen (TOMs) bereitgestellt werden. Eine Budgetkürzung bei der IT-Sicherheit kann im Schadensfall als Organisationsverschulden ausgelegt werden.
2. **Kontrollfunktion:** Er muss verifizieren, dass die finanzierten Massnahmen (z. B. Backups, IKS-Kontrollen im Zahlungsverkehr) tatsächlich wirksam sind und nicht nur auf dem Papier existieren. «Blindes Vertrauen» in die IT-Abteilung genügt den gesetzlichen Sorgfaltspflichten nicht mehr.

Fazit: Resilienz als Wettbewerbsvorteil

Die Bedrohungslage wird sich durch den Einsatz von künstlicher Intelligenz auf der Angreiferseite weiter verschärfen. Absolute Sicherheit ist eine Illusion. Das Ziel moderner Finanzführung muss daher die **Operationelle Resilienz** sein: Die Fähigkeit, Angriffe zu absorbieren, den geschäftskritischen Betrieb aufrechtzuerhalten und schnell zum Normalzustand zurückzukehren.

Unternehmen, die Datenschutz und Sicherheit proaktiv in ihre Prozesse integrieren, technische und organisatorische Massnahmen wie auch eine Awarenesskultur schaffen, minimieren nicht nur ihre Risiken, sondern generieren einen Wettbewerbsvorteil. In einem Markt, der auf Vertrauen basiert, wird der Nachweis robuster Sicherheitsstandards und eines verantwortungsvollen Umgangs mit Daten zu einem Qualitätsmerkmal, das Kunden, Investoren und Aufsichtsbehörden honorieren.

Digitale Transformation und Organisation



Markus Speck

7.1 Ambition und Zielbild

Digitale Transformationsprojekte haben ihren Ursprung entweder in einer klaren, top-down definierten Zielsetzung (z.B. bei Grossunternehmen) oder ergeben sich aus einem Leidensdruck im operativen Alltag (typisch bei KMU). In beiden Fällen sollten sich die Fach- und Führungskräfte zu Beginn einer Transformation genügend Zeit nehmen, eine Ambition mit Zielbild zu beschreiben.

«Think big, start small»

Think big:

- Definieren Sie den Zustand oder die Wirkung, welche das Transformationsprojekt in beispielsweise 12 – 24 Monaten erreicht haben wird.
- Beschreiben Sie die Verbesserungen, welche innerhalb des definierten Zeitraums realisiert sein werden.
- Halten Sie fest, welche Fach- und Führungskräfte die Transformation gestalten sollen und wie die Organisation von der Transformation profitieren wird.
- Beschreiben Sie, welche Effizienzgewinne erzielt werden sollen und wie diese für die Organisation mess- und spürbar sein werden.

Start small:

- Entscheiden Sie, mit welchem Thema begonnen werden und welches Anliegen als erstes gelöst sein soll.
- Erstellen Sie eine Grobplanung für weitere Arbeitspakete, welche bis zur Erreichung des Zielzustandes abgewickelt werden sollen.
- Legen Sie fest, in welchen zeitlichen Zyklen die Arbeitspakete realisiert werden sollen. Empfehlung: zeitliche Zyklen zwischen 1 bis 3 Monaten für die erfolgreiche Realisierung eines Arbeitspakets.

Wenn Ambition und Zielbild definiert sind, wird eine «Reiseplanung» erstellt. Dies entspricht einer professionellen Projektplanung mit Projektsponsoren (z.B. GL-Mitglied), Projektleitung und Projektmitarbeitenden. Die Projektabwicklung sollte in agilen Zyklen vorgesehen werden, da rasche und stetige Projekterfolge entscheidend für die Zielerreichung sind.

Transformationsprojekte, welche aus Leidensdruck angestossen werden

Bei KMU lösen oft Alltagsprobleme einen Handlungsbedarf aus. Beispiele: «Wir müssen endlich den Kreditorenablauf automatisieren» oder «Unser Monatsreporting muss aussagekräftiger gestaltet werden und rascher vorliegen». Selbstverständlich kann man solche Anliegen individuell lösen. Dabei vergibt man aber die Chance, sich mit wenig zeitlichem Einsatz ein Bild über weitere Verbesserungspotenziale zu erarbeiten. Beispiel: «Welche weiteren Optimierungen können wir in 6 Monaten erreichen? In 12 Monaten? In 18 Monaten?».

«Pain-Point-Themen» regen zum «grösser Denken» an, damit Ambition und Zielbild definiert werden können.

Case Study: Vom Pain Point «Preisverhandlungen» zum «Profitabilitäts-Portfolio»

«PETCARE AG» ist ein Handelsunternehmen für Tier-Arznei- und Tier-Nahrungsergänzungsprodukte. Der Kundestamm umfasst über 1000 Tierkliniken / Tierarztpraxen in der Schweiz. Das Anliegen von PETCARE AG war: «Wir wollen unser jährliches Abschlusspreis-System effizienter und wirksamer gestalten. Oft hat nur der Kunde eine WIN-Situation, aber wir nicht.»

Auf die Frage nach der Kundenprofitabilität lautete die Antwort: «Wir orientieren uns an den Umsätzen, die Kundenprofitabilität können wir leider nicht so einfach ableiten.» Es lag auf der Hand, dass jedoch genau die Transparenz nach Kundenprofitabilität zwingend notwendig ist, damit das Kernanliegen gelöst werden kann. Das Unternehmen arbeitete mit drei Systemen: ein Auftragsabwicklungssystem, ein CRM-System und ein Accounting-System. Es hätte Monate gebraucht, um die Werteflüsse in diesen Systemen so abzubilden, dass aus dem Accounting-System Auswertungen nach Kundenrentabilität verfügbar gewesen wären. Der Informationsbedarf musste jedoch in wenigen Wochen erfüllt werden. PETCARE AG entschied sich für eine Profitabilitätsanalyse mit Power-BI.

Kostenrechnerisch mussten folgende Vorbereitungsarbeiten getroffen werden:

- Detaillierte Deckungsbeitragsstruktur festlegen
- Datenquellen und Verfügbarkeit für jede Zeile der DB-Rechnung herleiten
- Prozesskostensätze für Auftragsabwicklung und Verkauf definieren
- Weitere Kosten- oder Zuschlagssätze bestimmen

Nach dem konzeptionellen Entscheid erfolgten die Datenmodellierung in Power-BI und der sukzessive Aufbau der Führungs-Cockpits. In weniger als sechs Wochen konnten bereits viele wertvolle Reports freigegeben werden. Die raschen Erfolge lösten weiteren Informationsbedarf aus. Die Geschäftsleitung stellte sich dank der gewonnenen Erkenntnisse verschiedene strategische Fragen zur eigenen Wettbewerbsstrategie. So wurde beispielsweise eine Portfolio-Sicht in Power-BI realisiert, welche die Umsätze und Deckungsbeiträge pro Kunde mit deren Potenzial (Grösse der Klinik) in Bezug setzte. Daraus konnte das Unternehmen konkretere Marketing- und Verkaufsförderungsmassnahmen ableiten.

7.2 Führungsverantwortung und Data-Governance

Data Governance umfasst Richtlinien, Prozesse und Verantwortlichkeiten für den Umgang mit Daten im Unternehmen. Sie stellt sicher, dass Daten korrekt, sicher und regelkonform genutzt werden. Ein gutes Governance-Framework unterstützt Compliance und minimiert Risiken.

Bei der Realisierung von ERP-Systemen oder ganzheitlichen Datenmanagement-Konzepten eines Unternehmens haben CFO-Funktion oder Accounting-Führungskräfte eine hohe Mitverantwortung. Diese umfasst die Erarbeitung, Festlegung und Durchsetzung von Anforderungen. Eine Beschränkung der Verantwortung auf die unmittelbaren Accounting-Systeme ist dabei nicht ausreichend. Vielmehr muss der Einfluss auf die Gestaltung aller Systeme und Datenquellen des Unternehmens geltend gemacht werden. Oft sind neben klassischen

ERP-Systemen weitere Anwendungen im Einsatz, beispielsweise CRM-Systeme⁴⁵, technische Konfigurationstools⁴⁶ oder Anwendungen für die operative Leistungserstellung. Eine zentral geführte Governance ist beispielsweise für die folgenden Entscheidungen notwendig:

- Festlegung, welche Systeme und Applikationen unter welchen Bedingungen im Unternehmen eingesetzt werden
- Durchgängigkeit und Einheitlichkeit von Stammdaten über alle im Unternehmen eingesetzten Systeme sicherstellen. Beispiele: einheitliche

⁴⁵CRM steht für «Customer-Relationship-Management». CRM-Systeme unterstützen die konsequente Ausrichtung eines Unternehmens auf seine Kunden und die Gestaltung der Kundenbeziehungsprozesse. Dies kann Aufgaben vom Kontakt-Management bis hin zur umfassenden Abbildung des Verkaufsprozesses umfassen.

⁴⁶Mit Konfigurationstools werden unter anderem die technischen Dimensionen eines Produktes oder einer Produkthanwendung verwaltet und bestimmt. Beispiel: das Konfigurationstool eines Fensterbauunternehmens verwaltet die geometrischen Ausmasse sowie die Art und Anzahl der eingesetzten Materialien für verschiedene Fenstertypen.

7 Digitale Transformation und Organisation

Artikelstammdaten, einheitliche Kontakt- und Adresstammdaten

- Mit dieser Durchgängigkeit verbunden: Vermeidung jeglicher Redundanzen, sei es bei Stammdaten, Prozessen oder bei der Verarbeitung von Daten
- Klare Vorgaben bezüglich Business-Intelligence-Lösungen, Robotics im Accounting, Automatisierungs-Routinen im Accounting
- Richtlinien zur Datenhaltung und Datenspeicherung (z. B. On-Premise, Cloud), bei Cloud-Lösungen: Vorgaben, wo (geografisch) diese Daten verwaltet werden dürfen
- Sicherstellung aller Anforderungen bezüglich Daten-Integrität und Daten-Sicherheit: Gerade beim Thema Cyber Security werden von einer CFO-Organisation klare Richtlinien und Absicherungsinstrumente erwartet, da jeder Cyber-Angriff erheblichen Schaden anrichten kann.

Case Study: Data-Governance in einem Familien-Konzern

Die «IMT Holding AG» ist die Muttergesellschaft einer international tätigen Unternehmensgruppe im hochwertigen Maschinenbau. Die Gruppe ist in drei Subholdings gegliedert, wobei zwei dieser Subholdings durch Akquisitionen erworben wurden. Bislang hatten die drei Subholdings eine weitgehende Freiheit in der Auswahl und Gestaltung ihrer IT-Systeme. So waren denn auch drei verschiedene ERP-Systeme im Einsatz. Der CFO der IMT Holding AG lancierte ein Projekt zur Harmonisierung des Datenmanagements innerhalb der Gruppe, insbesondere die Evaluation eines gruppenweit anwendbaren BI-Systems. In diesem Zusammenhang war es notwendig, allgemeingültige Richtlinien zum Datenmanagement durchzusetzen. Der CFO formulierte dazu verschiedene Anträge an den Verwaltungsrat der Gruppe, welche eingehend diskutiert und dann für alle Einheiten verbindlich verabschiedet wurden. Beispiele dieser Entscheidungen:

- Die Subholdings wurden verpflichtet, ihre IT-Strategie jährlich zu aktualisieren und dem Verwaltungsrat zur Genehmigung vorzulegen.
- Der Ersatz oder Ausbau von ERP-Lösungen blieb zwar weiterhin in der Autonomie der Subholdings, musste zur Nutzung von Synergien jedoch vorgängig mit den anderen Einheiten abgestimmt und vom Verwaltungsrat freigegeben werden.
- Den Subholdings wurde untersagt, eigene BI-Lösungen zu evaluieren und zu implementieren. Dieser Prozess wurde übergeordnet von der Holding an die Hand genommen, die Bedürfnisse der Subholdings wurden miteinbezogen.
- Es wurde eine für alle Einheiten verbindliche BI-Lösung ausgewählt und implementiert.
- In einem späteren Schritt wurden durch die Initiative des Holding-CFO weltweit einheitliche Normen und Überwachungssysteme für Cyber-Sicherheit eingeführt.

7.3 System- und Datenlandschaft

Die Praxis zeigt, dass in den wenigsten Unternehmen die «perfekte Welt» bezüglich eingesetzter Systeme und Datenquellen existiert. Gerade bei KMU ist die System- und Datenlandschaft oft über die Jahre gewachsen, je nach Art der Data-Governance mehr oder weniger kontrolliert. Typischerweise sind folgende Arten von Anwendungen im Einsatz (nicht abschliessend):

- Klassische ERP-Anwendungen mit umfassender Abdeckung aller betrieblichen Prozesse, Daten- und Werteflüsse (eher in Grossunternehmen)
- Accounting-Anwendungen mit Ausprägungen Finanzbuchhaltung, allenfalls mit integrierter Kostenrechnung
- Anwendungen zur Planung, Steuerung und Abwicklung der betrieblichen Leistungserstellung. Oft handelt es sich hier um branchenspezifische Lösungen.
- CRM-Systeme zur Planung, Steuerung und Auswertung aller Verkaufsaktivitäten
- Frontend-Tools zur Datenerfassung, meist cloudbasiert (z.B. Scanning-Tools, Auftragsdatenerfassung von Technikern auf Kundenaufträgen)
- Kollaborationstools zur effizienten Zusammenarbeit, i. d. R. cloudbasiert (z.B. Microsoft-Teams, Slack, Miro, Beekeeper etc.)

Im Weiteren verwenden viele Unternehmen öffentlich verfügbare Datenquellen (z.B. Wechselkurs-tabellen der ESTV) oder spezifische Datenquellen, welche von Branchenverbänden zur Verfügung gestellt werden.

Führungskräfte im Accounting sollten sich einen Überblick zu den im Unternehmen eingesetzten Systemen und Datenquellen verschaffen und im Sinn der Data Governance die Regeln zur Zulässigkeit und Nutzung mitgestalten. Dies aus den folgenden Gründen:

- Compliance und Cyber Security: Gerade bei cloudbasierten Lösungen muss sichergestellt werden, dass nur vom Unternehmen freigegebene Anwendungen im Einsatz sind.
- «Single Point of Truth» bei Stammdaten und Vermeidung von Daten-Redundanz: Hier muss bestimmt werden, welche Anwendung die Stammdaten «hoheitlich» verwaltet. Alle übrigen Anwendungen, welche die gleichen Daten nutzen, richten sich nach dem «Single Point of Truth» aus. Beispiele von nicht gewünschter Redundanz: mehrfach angelegte Kontaktinformationen für den gleichen Kunden, mehrfach hinterlegte Artikeldaten, «Free-style»-Kalkulation auf Excel ausserhalb der definierten Kalkulationsstrukturen.
- Abstimmung von Prozessen und Arbeitsabläufen: Auch hier geht es darum, dass Arbeitsprozesse grundsätzlich nur in einem führenden System abgewickelt werden und die weiteren Anwendungen entsprechend aus diesem System gespeist werden. Die mehrfache Abwicklung gleichartiger Abläufe muss vermieden werden.
- Welche Daten stehen wo und wofür zur Verfügung? Wenn die Effizienz in der Auswertung und dem Nutzen von Daten verbessert werden soll, muss geklärt werden, welches die geeignete Datenquelle zum Abrufen dieser Daten ist. Beispiel: Für die Gestaltung eines BI-Cockpits können verschiedene, aufeinander abgestimmte Datenquellen verknüpft werden.

7 Digitale Transformation und Organisation

Case Study: System- und Datenlandschaft als Voraussetzung für ein BI-Cockpit

Le Bon Fromage SA ist ein Unternehmen mit rund 100 Mitarbeitenden, das verschiedene Käsesorten bis zur Ausreifung pflegt (Affinage) und die Produkte anschliessend konsumfertig für den weltweiten Detailhandel portioniert und verpackt. Das Unternehmen verfügt über ein gut etabliertes ERP-System, welches Materialwirtschaft, Stücklisten- und Arbeitsplanverwaltung, Kalkulation sowie Rechnungswesen abbildet. Zudem sind ein so genanntes MES (Manufacturing Execution System) sowie eine Quality-Management-Software im Einsatz. Die notwendigen Auswertungen werden konventionell über Reports direkt aus diesen Systemen erstellt. Es besteht eine sehr grosse Anzahl von Reports, jedoch verlieren die Führungskräfte die Übersicht aufgrund dieser grossen Menge.

Aus diesem Grund entschied sich das Unternehmen, die ganze Berichterstattung strukturiert über ein Business-Intelligence-Cockpit neu zu organisieren. Dazu wurde unter Einbezug aller Führungskräfte eine einfache Illustration zur Systemlandschaft erstellt:

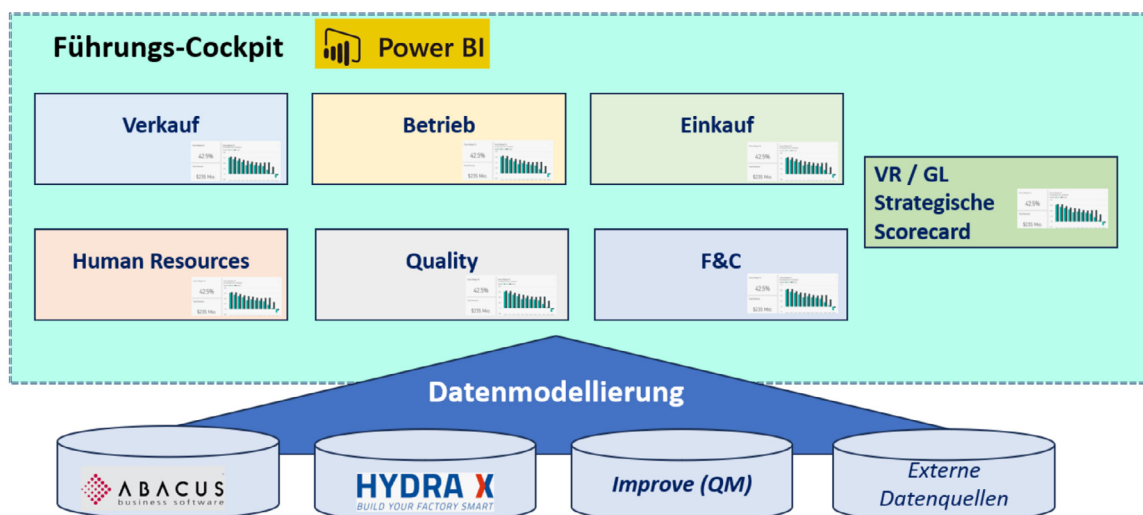


Abbildung 48: Illustration Systemlandschaft (eigene Darstellung)

Dank den bereits gut organisierten Daten in den einzelnen Systemen erfolgten die Datenmodellierung und die ersten Reports bereits innerhalb weniger Wochen. Den involvierten Führungskräften wurden dank dieser Darstellung die Zusammenhänge gut verständlich gemacht.

7.4 Projektmanagement und agile Methoden

Der klassische Projektansatz bei IT-Projekten erfolgte in der Vergangenheit oft nach dem so genannten «Wasserfall-Modell». Bei diesem Ansatz erfolgt die Projektabwicklung sequenziell und erstreckt sich über mehrere Monate oder meist über mehr als ein Jahr hinaus.

Wasserfall-Methode:

- Beschreibung der erwarteten Lösungen und Ergebnisse durch die Auftraggeber in einem Lastenheft, welches vom Projektteam in ein detailliertes Pflichtenheft übersetzt wird
- Detaillierte Definition aller Soll-Prozesse und Geschäftsvorfälle als Input für das Einrichten einer Testumgebung
- Einrichten und Testen aller Prozesse und Geschäftsvorfälle in der Testumgebung und sukzessiver Aufbau der produktiven Umgebung
- Einführung und umfassende Schulung der Prozesse und Geschäftsvorfälle für alle betroffenen Anwender in der Testumgebung. Nachschärfen der Anwendungen aufgrund der Schulungsfeedbacks

- Go-Live in der produktiven Umgebung. Nachbetreuung der Anwender und Bereinigung von aufgetauchten Mängeln

Die Erarbeitung von Lasten- und Pflichtenheften macht grundsätzlich Sinn. Bei einer sehr detaillierten Ausgestaltung besteht jedoch die Gefahr, dass der Anforderungskatalog zu umfangreich wird. Der Wunsch nach «umfassender Perfektion auf Anhieb» schafft wesentlich höhere Komplexität, bindet viele Ressourcen und kann zu Überlastungen im Projekt- ablauf führen. Aufgrund der langen Laufzeit liegen sicht- und greifbare Projektergebnisse erst spät vor. Nicht selten blähen zusätzliche Erkenntnisse und neue Wünsche während der Ausführung den Projektumfang zusätzlich auf.

Für Datenmanagement- und IT-Projekte wird ein agiler ⁴⁷ Projektansatz empfohlen. Ausgehend von einem Zielbild mit Projektdefinition (nicht zu detailliert) wird das Vorhaben in einzelne Arbeitspakete strukturiert. Diese werden in zeitlichen Zyklen von wenigen Wochen abgewickelt.

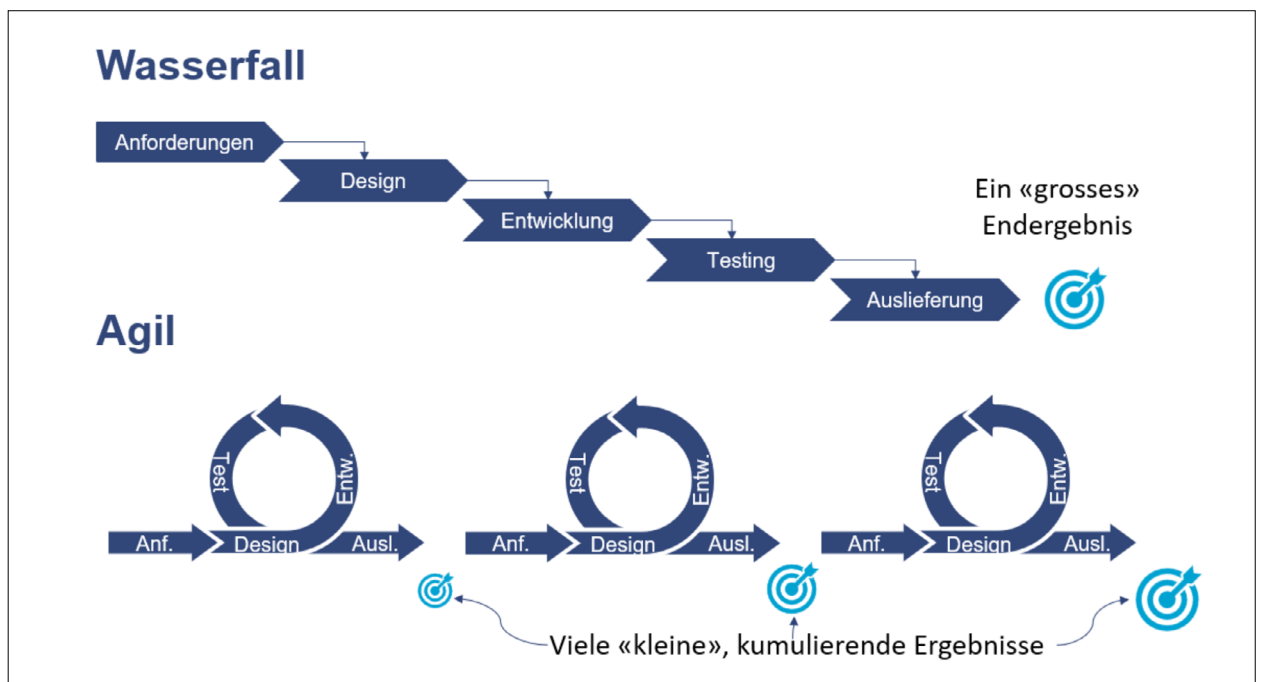


Abbildung 49: «Wasserfall» versus «Agil» (Quelle: datavision.ch)

⁴⁷ Agiles Projektmanagement ist ein Ansatz, bei dem Projekte in kurzen, überschaubaren Zyklen geplant und umgesetzt werden. Damit kann flexibel auf Veränderungen reagiert werden.

7 Digitale Transformation und Organisation

Bei der Methode «Agil» wird jeder Zyklus sauber mit Abnahme und Qualitätskontrolle abgeschlossen. Gewonnene Erkenntnisse aus einem abgeschlossenen Zyklus fließen in den nächsten Zyklus ein. Übersicht zum Vorgehen:

- Zielbild mit Projektdefinition erarbeiten und kommunizieren
- Betroffene Anwender so früh als möglich in die Arbeitspakete einbeziehen
- Pakete im Projekt priorisieren: Mit welchem Paket soll begonnen werden, welche Pakete bleiben im Backlog bzw. folgen in welcher Sequenz
- Inhalte und Vorgehen für das priorisierte Paket detailliert spezifizieren; Terminplan mit Aktivitäten und erwarteten Ergebnissen erstellen
- Inhalte jedes Arbeitspaketes innerhalb von 2 bis 3 Monaten bis zum gewünschten Ergebnis abwickeln
- Betroffene Anwender während der Abwicklungsphase eines jeden Pakets involvieren: Schulung,

Fine-Tuning der Anforderungen, Abnahme und Qualitätskontrolle

- «Lessons learned» und Feedback als Input für die folgenden Pakete verwenden

Der Vorteil des agilen Vorgehens liegt darin, dass viel rascher Erfolgserlebnisse vorliegen und dies einen kontinuierlichen Lern- und Verbesserungsprozess unter allen Beteiligten fördert. Auch agiles Vorgehen erfordert klare Verantwortungen und eine professionelle Projektorganisation.

- Projekt-Sponsoren und Steuerungsausschuss (Geschäftsleitungsebene)
- Leitung des Gesamtprojektes
- Verantwortliche für die Teilprojekte und Arbeitspakete
- Gesamt-Projektplanung mit Steuerung aller Arbeitspakete bis zum Projektabschluss
- Wirkungskontrolle nach jedem Arbeitspaket durch Anwender und Projektleitung

7 Digitale Transformation und Organisation

Case Study: Realisierung eines BI-Cockpit mit agilem Vorgehen

Das bereits vorgestellte Unternehmen Le Bon Fromage SA entschied sich zu einem agilen Vorgehen in der Realisierung des BI-Cockpits. Der CFO übernahm die Gesamt-Projektleitung, die Projektsteuerung wurde durch den CEO sowie die Geschäftsleitungsmitglieder sichergestellt. Als erstes Arbeitspaket wurde die Ausarbeitung des Verkaufs-Cockpits definiert. Zusammen mit dem Verkaufsleiter wurden die Anforderungen detailliert festgelegt. Hier eine stark gekürzte Fassung:

Le Bon Fromage SA - Absatz-, Umsatz- und Margenreporting

Was ist der Zweck eines solchen Reporting und wie soll es funktionieren?

- Das Reporting soll vom Leiter Verkauf, CEO und CFO aktiv und regelmässig genutzt werden.
- Weitere Anwender/innen im Verkauf werden in einem Berechtigungskonzept bestimmt.
- Das Reporting soll grundsätzlich in Echtzeit genutzt werden können (tagesaktuell bei Bedarf).

Welche Informationen sollen aus dem Reporting gewonnen werden?

- Absatzmengen pro Kunde / Artikel
- Umsatz pro Kunde / Artikel: in Originalwährung sowie in Buchungswährung (CHF)
- Deckungsbeitrag pro Kunde / Artikel (Wertschöpfung minus Affinage und Kosten VVP)

In welchen zeitlichen Dimensionen soll ausgewertet werden können?

- Ist laufendes Geschäftsjahr (einzelne Monate, Selektion von Monaten) und zwei Vorjahre
- Budget laufendes Jahr, rollierender Forecast 6 Quartale
- Analysen mit Vergleichen zu Vorjahr und Budget

Welche Datentabellen brauchen wir?

Artikelstamm, Kundenstamm, Artikelkalkulationen - Abacus
 Währungstabellen, Datumstabellen - Excel
 Bewegungsdaten Absatz/Umsatz - Abacus

Abbildung 50: Anforderungen Verkaufs-Cockpit (eigene Darstellung)

Auswertungsebene

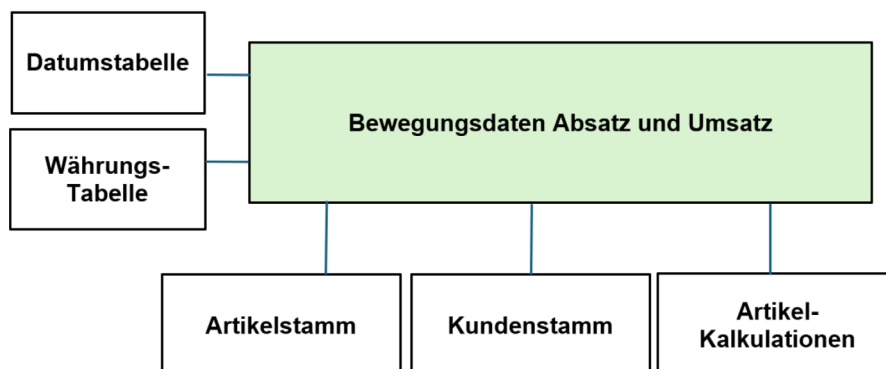


Abbildung 51: Anforderungen Verkaufs-Cockpit (eigene Darstellung)

Zusammen mit einem Power-BI-Experten wurden innerhalb von rund acht Wochen alle gewünschten Reports eingerichtet und vom Verkaufsleiter mit seinem Team getestet. Die Bereinigung aufgrund von Rückmeldungen nahm weniger als vier Wochen in Anspruch, sodass dieses Arbeitspaket in weniger als zwölf Wochen abgeschlossen und genutzt werden konnte.

7.5 Change Management als zwingender Erfolgsfaktor

Transformationsprojekte werden nur gelingen, wenn dem Change Management genügend Aufmerksamkeit und Zeit geschenkt werden. Veränderungen bei den Menschen innerhalb einer Organisation zu bewegen ist emotional wesentlich anstrengender als die technische Realisierung. Change Management kann wie ein Puzzle mit vier Bestandteilen dargestellt werden:



Abbildung 52: Change-Management-Puzzle (eigene Darstellung)

Change Management beginnt bei den obersten Führungskräften eines Unternehmens (Verwaltungsrat, Geschäftsleitung). Wenn diese Führungskräfte die Zielsetzungen einer Transformation nicht mittragen oder unzureichend verstehen, werden die Menschen der Organisation ambivalente Signale erhalten. Stellen Sie als erstes sicher, dass ein Transformationsprojekt von den obersten Führungskräften nicht nur verstanden, sondern überzeugt mitgetragen und kommuniziert wird.

Häufige Versäumnisse im Change Management und wie diese vermieden werden können:

Versäumnis / Mangel	Wie vermeiden
Zielsetzung der Transformation ist unklar, Konsequenzen werden verschwiegen	· Klares Zielbild und klare Reiseplanung kommunizieren (z.B. wo wollen wir in 24 Monaten sein, in welchen Schritten wollen wir dorthin gelangen) · Konsequenzen der Veränderungen ehrlich und transparent kommunizieren
Führungskräfte sind nur oberflächlich informiert	· Führungskräfte verbindlich einbeziehen, Workshop zum Verständnis von Zielen und dem Weg zum Ziel · Erwartungen an die Führungskräfte klar definieren
Mitarbeitende / Anwender werden zu spät oder zu oberflächlich einbezogen.	· Projekt mit Zielbild und Reiseplanung so früh als möglich vorstellen · Aufzeigen, wer in welcher Phase mit welchen Aufgaben im Projekt mitarbeiten wird · Ausblick zeigen: welche Rolleninhalte sich wie ändern werden, welche Systeme und Strukturen künftig wie genutzt werden
Zu kurzes oder zu wenig fundiertes Training in der neuen Anwendung	· Einführungs- und Trainingsplan für die Anwender frühzeitig erstellen · Anliegen der Anwender früh abholen und integrieren · Erwartungen und Trainingsziele definieren (was muss nach dem Training beherrscht werden) · Initialtrainings und «Wiederholungskurse» durchführen · Feedback der Anwender für kontinuierliche Verbesserung nutzen
Unausgesprochene Ängste der Mitarbeitenden	· Ängste und Unbehagen bei den Mitarbeitenden frühzeitig offen ansprechen und behandeln · Ehrlichkeit bezüglich künftiger Aufgaben und Rollen der Mitarbeitenden
Widerstände von Mitarbeitenden, «Vermeidungshaltung»	· Widerstände und Hindernisse zum Ziel offen ansprechen und gemeinsam lösen · Vereinbarte Lösung festhalten und Verbindlichkeit einfordern

7 Digitale Transformation und Organisation

Umgang mit Widerständen und Hindernissen

Widerstände werden als unangenehm wahrgenommen, können beim richtigen Umgang aber positive Energie freisetzen. Grundsätzlich sollten skeptische Stimmen ernst genommen werden. Mit einem gemeinsamen Vorgehen können Widerstände und Hindernisse konstruktiv beseitigt werden. In einem Workshop führen drei Fragen oder Arbeitsschritte zum Ziel:

1. **«Wo stehen wir heute, wo wollen wir bei Erreichung der Projektziele stehen?»** Mit dieser Frage wird eine allenfalls emotionale Ausgangslage versachlicht.

2. **«Welche Hindernisse müssen wir auf dem Weg vom Heute zum Ziel überwinden?»** Konkrete Anliegen der Skepsis werden hier von allen Beteiligten offengelegt.
3. **«Wie gehen wir vor, damit die Hindernisse überwunden werden?»** Nun sind die konstruktiven Beiträge aller Beteiligten gefordert. Dieser Punkt wird erst abgeschlossen, wenn die Lösungen auf dem Tisch liegen.

Zum Schluss noch dieses Zitat:

«Wer Lösungen will, sucht Wege. Wer keine will, sucht Gründe».

Case Study: Change Management in einem Transformationsprojekt

Die «Holzbau AG» ist ein Unternehmen mit rund 150 Mitarbeitenden und realisiert Projekte für hochwertige Bauten im gewerblichen Bereich sowie im Wohnbau. Das Unternehmen hat vor Jahren ein auf ihre Kernprozesse zugeschnittenes Auftragsbearbeitungssystem entwickeln lassen und führt neben diesem System das Rechnungswesen in Abacus, ein branchenspezifisches Konfigurationstool sowie eine separate CRM-Anwendung. Der neu in die Verantwortung gekommene CFO erkannte, dass innerhalb der kommenden drei Jahre das Auftragsbearbeitungssystem ersetzt werden muss und eine stärkere Integration mit den übrigen Systemen notwendig ist. So war beispielsweise die Nachkalkulation der Kundenaufträge unzureichend, und die Kostenrechnung innerhalb des Rechnungswesens war nur sehr rudimentär ausgeprägt.

Als ersten Schritt organisierte der CFO einen gemeinsamen Workshop mit allen Betreuungspersonen der verschiedenen Systeme. Dabei wurden die Ausgangslage mit ihren Stärken und Schwächen gewürdigt und ein grobes Zielbild entwickelt. Während diesem Workshop wurde beispielsweise festgehalten, dass Kontakt- und Adressdaten unabhängig voneinander in drei verschiedenen Systemen gepflegt wurden und aus diesem Grund eine hohe Redundanz vorhanden war. Die Bereinigung dieser Situation wurde als ein erstes Arbeitspaket priorisiert. Als weiterer Schwerpunkt wurde die Überarbeitung der Prozesse für Beschaffung, Wareneingang und Kontrolle/Freigabe von Lieferantenrechnungen festgelegt. Diese Erkenntnisse behandelte der CFO anschliessend innerhalb der Geschäftsleitung.

Es zeigte sich schnell, dass eine rein informative Behandlung innerhalb der Geschäftsleitung nicht ausreichend war. So entschied sich das Gremium, zusammen mit den Systemverantwortlichen einen Gestaltungsworkshop von einem ganzen Tag zu organisieren. Dabei wurden unter anderem sämtliche unterschiedlichen Beschaffungsprozesse beurteilt und die Anforderungen an deren Sollzustand definiert. Vorbehalte und mögliche Hindernisse wurden offen angesprochen und gemeinsam beseitigt. Am Schluss des Workshops lagen eine klare Zielsetzung und Marschrichtung vor, die dann verschiedenen Fach- und Führungskräften erläutert wurden. Die Arbeiten im Workshop legten auch Schwachstellen in den Abläufen sowie fehlende IKS-Routinen offen. Diese Erkenntnisse wurden bei der Definition von Zielen und Vorgehensweisen berücksichtigt.

Die klare und verbindliche Haltung der Geschäftsleitungsmitglieder übertrug sich auch auf die in der Projektumsetzung involvierten Personen.

7.6 Zusammenfassung

Idealerweise steht am Beginn digitaler Transformationsprojekte ein Zielbild mit einer klaren Roadmap für die Umsetzung. Das Zielbild wird vom CFO und den Führungskräften im Accounting massgeblich mitgestaltet. Dem Denken in grossen Dimensionen («Think big») folgen sinnvoll portionierte Arbeitspakete. Es ist empfehlenswert, mit den dringendsten Anliegen des Geschäftsalltags zu beginnen («Start small»), um dann sukzessive die weiteren Projektschritte bis zum Zielzustand durchzuführen.

In einer CFO-Funktion ist es zwingend, die Data Governance eines Unternehmens aktiv mitzugestalten. Dies geht weit über die Auswahl von Software-Anwendungen hinaus. Vielmehr sind es Themen wie Datenintegrität, Vermeidung von Prozess- und Datenredundanzen, Cyber-Security-Richtlinien, Entscheide über Datenhaltung (On-Premise, Cloud, welche Clouds) etc. Im gleichen Zusammenhang können auch die Kennnisse und die Mitentscheidungen über eingesetzte Datenquellen genannt werden. CFOs und Führungskräfte im Accounting stellen einen permanenten Überblick zu den internen und externen Datenquellen sicher. Zusammen mit den anderen Führungsorganen legen sie die Regeln zur Einbindung und Nutzung der Datenquellen fest.

Bei IT- oder Datenmanagement-Projekten können entweder die klassische «Wasserfall-Methode» oder ein agiles Vorgehen gewählt werden. Die Wasserfall-Methode startet mit detaillierten Spe-

zifikationen der erwarteten End-Ergebnisse. Dann folgen in sequenzieller Folge Gestaltungs- und Testphase, Implementierung, Schulung und Go-Live. Beim agilen Vorgehen wird der Ziel-Zustand genügend genau beschrieben und die Abwicklung danach in Arbeitspakete gegliedert und priorisiert. Die Arbeitspakete werden gemäss Prioritätenliste in Zyklen von wenigen Wochen bis zum erwarteten Ergebnis abgewickelt. Bei der Wasserfall-Methode liegt das grosse End-Ergebnis erst am Schluss vor, beim agilen Vorgehen werden nach jedem Zyklus die gewünschten Ergebnisse erreicht.

Transformationsprojekte gelingen dann, wenn dem Change Management genügend Zeit und Ressourcen zur Verfügung stehen. Durch eine Transformation verändern sich Systeme und Strukturen sowie die Art und Weise, wie diese genutzt werden sollen. Die Rolleninhalte von Mitarbeitenden und Führungskräften müssen in der Regel an diese Änderungen angepasst werden. Dies muss klar und transparent kommuniziert werden. Die Befähigung der betroffenen Fachleute muss geplant und sichergestellt werden. Letztendlich muss jede von der Veränderung betroffene Person entscheiden, ob sie auf die Reise der Veränderung mitgehen will oder nicht.

Glossar⁴⁸

API (Application Programming Interface)

Programmierschnittstelle, die es verschiedenen Softwareanwendungen ermöglicht, miteinander zu kommunizieren und Daten auszutauschen. APIs arbeiten im Hintergrund und sind stabiler als GUI-basierte Automatisierung, da sie von Änderungen der Benutzeroberfläche unberührt bleiben.

Attended Automation (Beaufsichtigte Automatisierung)

RPA-Variante, bei der Software-Roboter als digitale Assistenten auf dem Arbeitsplatzrechner eines Mitarbeitenden arbeiten. Der Bot wird durch Benutzerinteraktion ausgelöst und übernimmt Teilaufgaben, während der Mensch die Kontrolle behält.

Batch-Verarbeitung (Stapelverarbeitung)

Datenverarbeitungsmethode, bei der Daten in geplanten Intervallen gesammelt und gemeinsam verarbeitet werden. Kostengünstig für grosse Datenmengen mit niedriger Dringlichkeit. Typische Anwendungen: Monatsabschlüsse, Gehaltsabrechnungen, nächtliche Datenaktualisierungen.

Bot (Software Roboter)

Softwareprogramm, das automatisiert Aufgaben ausführt, die ansonsten von Menschen erledigt werden müssten. Im RPA-Kontext keine physischen Roboter, sondern rein digitale Entitäten, die auf Computersystemen laufen und dort Aufgaben wie Dateneingabe, Systemabfragen oder Dokumentenverarbeitung übernehmen.

BPMS (Business Process Management System)

Software zur Modellierung, Ausführung und Optimierung von Geschäftsprozessen. Im Gegensatz zu RPA erfordert klassische BPMS-Automatisierung oft tiefgreifende Systemanpassungen und neue Schnittstellen, was höhere Implementierungskosten und längere Projektlaufzeiten bedeutet.

Business Intelligence

Sammelbegriff für Methoden, Prozesse und Technologien zur systematischen Sammlung, Aufbereitung, Analyse und Darstellung von Unternehmensdaten, mit dem Ziel, Entscheidungen zu unterstützen und Steuerungsinformationen bereitzustellen.

Change Management

Systematischer Ansatz zur Begleitung von Veränderungsprozessen in Organisationen.

Cloud-Automatisierung

Automatisierung von Geschäftsprozessen über cloudbasierte Plattformen wie Zapier, Make oder Power Automate. Nutzt typischerweise APIs statt GUI-Simulation, ist daher stabiler als klassisches RPA, setzt aber voraus, dass die zu verbindenden Anwendungen entsprechende Schnittstellen anbieten.

Compliance

Einhaltung gesetzlicher Bestimmungen, regulatorischer Vorgaben und interner Richtlinien. Für Schweizer Unternehmen relevant: DSG (Schweiz), DSGVO (EU-Bezug), FINMA-Rundschreiben (Finanzsektor). Bei Cloud-Automatisierung besonders wichtig: Serverstandort, Datenverschlüsselung, Auftragsverarbeitungsverträge.

CRM (Customer Relationship Management)

Eine CRM-Applikation ist eine Softwarelösung zur zentralen Erfassung, Verwaltung und Analyse aller Kundeninteraktionen und -daten, die Unternehmen dabei unterstützt, ihre Prozesse in den Bereichen Vertrieb, Marketing und Service zu optimieren und die Kundenbeziehungen nachhaltig zu stärken.

Data Governance

Data Governance umfasst Richtlinien, Prozesse und Verantwortlichkeiten für den Umgang mit Daten im Unternehmen. Sie stellt sicher, dass Daten korrekt, sicher und regelkonform genutzt werden. Ein gutes Governance-Framework unterstützt Compliance und minimiert Risiken.

⁴⁸In Anlehnung an Wikipedia, Copilot, eigene Definitionen

Glossar

Data Lake

Speicherarchitektur, die grosse Mengen von Rohdaten in ihrem ursprünglichen Format speichert – sowohl strukturierte als auch unstrukturierte Daten. Bildet die Basis für ELT-Ansätze, bei denen die Transformation erst bei Bedarf erfolgt.

Data Mesh

Data Mesh ist ein dezentraler Architekturansatz, der die Verantwortung für Daten von einem zentralen Team auf fachliche Domänen-Teams verlagert und Daten dabei als Produkte betrachtet, die im Unternehmen bereitgestellt werden.

Data Ownership

Data Ownership definiert die organisatorische Eigentümerschaft und Verantwortung für spezifische Datenbestände, einschliesslich der Zuständigkeit für deren Qualität, Sicherheit und konformen Nutzung über den gesamten Lebenszyklus.

Data Warehouse

Zentrales Datenlager für strukturierte, bereits transformierte Daten, optimiert für Analysen und Reporting. Typisches Zielsystem für ETL-Prozesse, in dem Daten in einem einheitlichen, analysereifen Format vorliegen.

Datenqualität

Mass für Vollständigkeit, Genauigkeit, Konsistenz und Aktualität von Daten.

Delta-Table

Eine Delta-Table ist ein erweitertes Speicherformat für Data Lakes, das Transaktionssicherheit (ACID), Versionierung und eine zuverlässige Datenverwaltung auf Basis von Parquet-Dateien ermöglicht.

DSGVO (Datenschutz-Grundverordnung)

EU-Verordnung zum Schutz personenbezogener Daten, gültig seit 2018. Für Schweizer Unternehmen relevant bei Geschäftsbeziehungen mit EU-Kunden oder -Partnern. Unterschiede zum Schweizer DSG: 72-Stunden-Meldefrist (statt «so rasch wie möglich»), Bussen bis 4% des Jahresumsatzes, Cookie-Banner-Pflicht.

EDÖB (Eidg. Datenschutz- und Öffentlichkeitsbeauftragter)

Schweizer Aufsichtsbehörde für Datenschutz. Zuständig für die Überwachung der Einhaltung des DSG. Bei Datensicherheitsverletzungen mit hohem Risiko ist eine Meldung an den EDÖB «so rasch wie möglich» erforderlich.

Echtzeitverarbeitung (Real-time Processing)

Datenverarbeitungsmethode, bei der Daten unmittelbar nach Eingang verarbeitet werden. Unerlässlich für zeitkritische Anwendungen wie Betrugserkennung, Börsenhandel oder Gesundheitsmonitoring. Höhere Infrastrukturkosten als Batch-Verarbeitung.

Employee Self-Service (ESS)

Digitale Portale, über die Mitarbeitende administrative Aufgaben selbstständig erledigen können – etwa Urlaubsanträge, Zeiterfassung, Stammdatenänderungen oder Speseneinreichungen. Entlastet die Personalabteilung und erhöht die Mitarbeiterzufriedenheit durch schnellere Bearbeitung.

ERP (Enterprise Resource Planning)

ERP-Systeme sind integrierte Softwarelösungen, die alle zentralen Geschäftsprozesse wie Finanzen, Einkauf, Produktion und Personalwesen in einer gemeinsamen Datenbank abbilden. Sie ermöglichen eine durchgängige Datenverarbeitung und verbessern die Transparenz sowie Effizienz im Unternehmen.

ETL (Extract, Transform, Load)

Prozess der Datenintegration, bei dem Daten aus verschiedenen Quellen extrahiert, in ein einheitliches Format transformiert (bereinigt und aufbereitet) und anschliessend in ein Zielsystem – wie etwa ein Data Warehouse – geladen werden, um sie für Analysen verfügbar zu machen.

GUI (Graphical User Interface)

Grafische Benutzeroberfläche einer Software. RPA-Bots können über die GUI mit Anwendungen interagieren, indem sie Mausklicks, Tastatureingaben und Bilschirmelemente simulieren. GUI-Automatisierung ist anfälliger für Änderungen als API-basierte Automatisierung.

Glossar

Human-in-the-Loop

Automatisierungsansatz, bei dem Menschen an definierten Kontrollpunkten eingreifen, Entscheidungen treffen oder Ausnahmen behandeln. Kombiniert maschinelle Effizienz mit menschlicher Urteilsfähigkeit. Insbesondere bei sensiblen Entscheidungen (z.B. Kreditvergabe) oft regulatorisch erforderlich.

Hyperautomation (Intelligent Process Automation)

Verschmelzung von RPA mit künstlicher Intelligenz, Machine Learning und Process Mining zu lernfähigen End-to-End-Automatisierungen. Erweitert klassisches RPA um Fähigkeiten wie intelligente Dokumentenverarbeitung, Natural Language Processing und kontextuelle Entscheidungsfindung.

KPI (Key Performance Indicator)

Kennzahl zur Messung des Erfolgs von Automatisierungsprojekten. Relevante KPIs: eingesparte Bearbeitungszeit pro Vorgang, Reduktion der Fehlerquote, Verkürzung von Durchlaufzeiten, Kosteneinsparungen, Return on Investment (ROI).

Legacy-System (Altsystem)

Ältere IT-Systeme, die noch im Einsatz sind, aber keine modernen Schnittstellen (APIs) bieten. RPA ist oft die einzige Möglichkeit, solche Systeme in automatisierte Workflows einzubinden, da die Bots über die Benutzeroberfläche interagieren können.

Low-Code / No-Code

Entwicklungsansatz, bei dem Anwendungen mit minimaler oder ohne Programmierung erstellt werden. RPA-Plattformen und Cloud-Automatisierungstools bieten grafische Oberflächen mit Drag-and-Drop, die auch Mitarbeitenden ohne Programmierkenntnisse die Erstellung von Automatisierungen ermöglichen.

Machine Learning

Machine Learning als Teilbereich der Künstlichen Intelligenz (KI) umfasst Verfahren, bei denen Algorithmen durch das Erkennen von Mustern in Daten selbstständig lernen, Probleme zu lösen oder Vorschläge zu treffen.

Medallion-Architektur

Die Medallion-Architektur ist ein Datenmodellierungskonzept für Data Lakehouses, das Daten logisch in drei Qualitätsstufen (Bronze für Rohdaten, Silber für bereinigte Daten, Gold für aggregierte Geschäftsdaten) unterteilt.

OCR (Optical Character Recognition)

Technologie zur automatischen Texterkennung in gescannten Dokumenten oder Bildern. Ermöglicht RPA-Bots die Verarbeitung von Papierdokumenten, PDFs und anderen nicht-digitalen Quellen. Moderne KI-gestützte OCR erreicht Erkennungsraten von über 99 %.

On-Premise

On-Premise bezeichnet ein Lizenz- und Nutzungsmodell, bei dem Software auf den eigenen Servern und der IT-Infrastruktur des Unternehmens vor Ort installiert und betrieben wird, anstatt in einer externen Cloud.

PoC (Proof of Concept)

Machbarkeitsstudie zur Validierung eines Automatisierungskonzepts vor der vollständigen Implementierung. Sollte einen repräsentativen, aber überschaubaren Prozess umfassen, der innerhalb weniger Wochen automatisiert werden kann.

PoV (Proof of Value)

Erweiterung des PoC, die nicht nur die technische Machbarkeit, sondern den konkreten Geschäftsnutzen nachweist. Messbare Kriterien: eingesparte Zeit, reduzierte Fehlerquote, verkürzte Durchlaufzeiten. Bildet die Grundlage für Investitionsentscheidungen.

Privacy by Design / by Default

Prinzip des nach dem seit 1. September 2023 geltenden revidierten Datenschutzgesetzes, DSGVO, und der DSGVO: Datenschutz muss bereits bei der Entwicklung und Konzeption von Systemen berücksichtigt werden (by Design) und die datenschutzfreundlichste Einstellung muss Standard sein (by Default).

Glossar

Process Mining

Analysemethode zur Rekonstruktion und Visualisierung von Geschäftsprozessen aus System-Logs und Transaktionsdaten. Identifiziert Engpässe, Abweichungen und Automatisierungspotenziale. Oft als Vorstufe zu RPA-Implementierungen eingesetzt.

Python

Python ist eine vielseitige, gut lesbare Programmiersprache, die aufgrund ihrer umfangreichen Bibliotheken besonders häufig in den Bereichen Data Science, künstliche Intelligenz und Automatisierung eingesetzt wird.

Quick Win

Überschaubares Automatisierungsprojekt mit schnell sichtbarem Nutzen und geringem Risiko. Dient dem Aufbau von Erfahrung, der Demonstration von Mehrwert und der Schaffung von Akzeptanz für weitere Automatisierungsvorhaben. Typisch: 2–4 Wochen Implementierungszeit.

Real-Time Daten

Real-Time Daten (Echtzeitdaten) sind Informationen, die unmittelbar im Moment ihrer Entstehung erfasst, verarbeitet und bereitgestellt werden, um eine verzögerungsfreie Analyse oder Reaktion zu ermöglichen.

RPA (Robotic Process Automation)

Technologie zur Prozessautomatisierung, bei der Software-Roboter («Bots») repetitive, regelbasierte Aufgaben übernehmen, indem sie menschliche Interaktionen auf der Benutzeroberfläche von Anwendungen nachahmen.

SaaS (Software-as-a-Service)

Software-as-a-Service (SaaS) ist ein cloudbasiertes Bereitstellungsmodell, bei dem Softwareanwendungen von einem externen Anbieter gehostet und dem Nutzer über das Internet – meist auf Abonnementbasis – zur Verfügung gestellt werden, ohne dass eine lokale Installation oder Wartung auf dem eigenen Rechner erforderlich ist.

SQL (Structured Query Language)

SQL ist eine standardisierte Datenbanksprache zur Definition, Abfrage und Manipulation von strukturierten Daten in relationalen Datenbanksystemen.

Stream-Processing

Kontinuierliche Verarbeitung von Datenströmen in Echtzeit, ohne diese zwischenspeichern. Typische Anwendungen: IoT-Sensordaten, Social-Media-Feeds, Log-Analyse. Ermöglicht unmittelbare Reaktionen auf eingehende Ereignisse.

Trigger

Auslösendes Ereignis, das eine Automatisierung startet. Kann zeitbasiert (z.B. täglich um 8 Uhr), ereignisbasiert (z.B. neue E-Mail) oder manuell sein. Grundprinzip aller Cloud-Automatisierungen: Trigger + Aktion(en) = Automation.

Unattended Automation (Unbeaufsichtigte Automatisierung)

RPA-Variante, bei der Bots vollkommen selbständig ohne menschliche Interaktion arbeiten. Nach Aktivierung durch einen Trigger führen sie transaktionsbasierte Aktivitäten aus. Typisch für Back-Office-Prozesse mit hohem Volumen wie Massen-datenverarbeitung oder nächtliche Batch-Jobs.

Vendor Lock-in (Anbieterabhängigkeit)

Risiko bei Cloud-Diensten: Je tiefer die Integration, desto schwieriger der Wechsel zu einem anderen Anbieter. Gegenmassnahmen: Export-Funktionen prüfen, offene Standards bevorzugen, Exit-Strategie von Beginn an einplanen, Dokumentation aller Workflows.

Workflow

Definierte Abfolge von Arbeitsschritten zur Erledigung einer Aufgabe. In Automatisierungsplattformen werden Workflows visuell als Flussdiagramme modelliert, bestehend aus Triggern, Aktionen, Bedingungen und Verzweigungen.

Literatur- und Quellenverzeichnis

Abacus (2025), <https://www.bsi-partner.ch/abacus/> (Abruf am 27.06.2026)

Abbyy (2024): State of Intelligent Automation Report – GenAI Confessions 2025.

aR solutions (2025): Revisionssichere Archivierung – Was bedeutet das? URL: <https://www.ar-solutions.ch/blog/digitale-archivierung-was-bedeutet-revisionssicher> (Abruf: 27.06.2026).

Bexio (2029, Bexio API, <https://docs.bexio.com/#tag/Reports/operation/ListJournalEntries> (Abruf 27.06.2026).

Bundesamt für Cybersicherheit BACS (2024): Wochenrückblick 14 – Online meeting with deepfake boss: CEO fraud 2.0. URL: https://www.ncsc.admin.ch/ncsc/en/home/aktuell/im-fokus/2024/wochenrueckblick_14.html (Abruf: 01.07.2026).

Bundesamt für Cybersicherheit BACS (2025): Halbjahresbericht 2024/2 – Cyberbedrohungslage Schweiz. URL: <https://www.ncsc.admin.ch/ncsc/de/home/dokumentation/berichte/lageberichte/halbjahresbericht-2024-2.html> (Abruf: 27.06.2026).

Bundesversammlung der Schweizerischen Eidgenossenschaft (2020a): Bundesgesetz über den Datenschutz (DSG) vom 25. September 2020, SR 235.1, in Kraft seit 01.09.2023. URL: <https://www.fedlex.admin.ch/eli/cc/2022/491/de> (Abruf: 27.06.2026).

Bundesversammlung der Schweizerischen Eidgenossenschaft (2020b): Bundesgesetz über die Informationssicherheit beim Bund (Informationssicherheitsgesetz, ISG) vom 18. Dezember 2020, SR 128, Art. 74a (Meldepflicht in Kraft seit 1. April 2025). URL: <https://www.fedlex.admin.ch/eli/oc/2022/232/de> (Abruf: 01.07.2026).

DAMA International (2017): DAMA-DMBOK.[1][2] Data Management Body of Knowledge. 2. Auflage. Basking Ridge: Technics Publications.

Databricks (2024): ETL und ELT – Eine ausführliche Betrachtung zweier Datenverarbeitungsansätze. <https://www.databricks.com/discover/etl>

Databricks (2025): Delta Lake Guide. URL: <https://docs.databricks.com/delta/index.html> (Abruf: 24.12.2025).

Databricks (2025): What is the Medallion Architecture? URL: <https://www.databricks.com/blog/what-is-medallion-architecture> (Abruf: 27.06.2026).

DataVision (2026), Automatisierung und Analytics im Accounting & Controlling, <https://datavision.ch/> (Abruf 27.06.2026).

Dehghani, Zhamak (2022): Data Mesh. Delivering Data-Driven Value at Scale. Sebastopol: O'Reilly.

economiesuisse (2021): Datenschutz – eine Übersicht zum neuen Gesetz, 19.05.2021. URL: <https://www.economiesuisse.ch/de/artikel/datenschutz-eine-uebersicht-zum-neuen-gesetz-0> (Abruf: 27.06.2026).

EDÖB: Erläuterungen zum Datenschutzgesetz. <https://www.edoeb.admin.ch/de>

Eidgenössische Finanzmarktaufsicht FINMA (2022): Rundschreiben 2023/1 «Operationelle Risiken und Resilienz – Banken», in Kraft seit 01.01.2023. URL: <https://www.finma.ch/de/~media/finma/dokumente/dokumentencenter/myfinma/rundschreiben/finmars-2023-01-20221207.pdf> (Abruf: 27.06.2026).

Eidgenössische Finanzmarktaufsicht FINMA (2024): Aufsichtsmitteilung 03/2024 vom 7. Juni 2024 – Meldepflichten bei Cyberangriffen nach Art. 29 Abs. 2 FINMAG (EHP, 24 h / 72 h). URL: <https://www.finma.ch/news/2024/06/20240607-mm-am-cyberisiken/> (Abruf: 12.05.2026).

FINMA: Rundschreiben 2018/3 (Outsourcing) https://www.finma.ch/de/~media/finma/dokumente/dokumentencenter/myfinma/rundschreiben/finmars-2018-03-01012021_de.pdf?sc_lang=de&hash=0DE-C4E84D9B2C22E00AD58039A6FBCFE (Abruf: 01.07.2026) und 2023/1 (Operationelle Risiken) https://www.finma.ch/de/~media/finma/dokumente/dokumentencenter/myfinma/rundschreiben/finmars-2023-01-20221207.pdf?sc_lang=de&hash=3DA82629BEB-D5388845AB8FD93121801 (Abruf: 01.07.2026)

Gabler (2025), Gabler Wirtschaftslexikon. <https://wirtschaftslexikon.gabler.de/definition/vuca-119684> (Abruf 127.06.2026)

Gartner (2024): Magic Quadrant for Robotic Process Automation.

Gartner (2025): Magic Quadrant for Analytics and Business Intelligence Platforms

Gartner (2025): Finance Analytics: Mature Data Management and Reporting Capabilities. <https://www.gartner.com/en/finance/topics/finance-analytics> (Abruf 27.06.2026)

Literatur- und Quellenverzeichnis

Harrison, Guy, (2015), Next Generation Databases. New York: Apress.

Haufe Akademie (2025): Change Management – Definition, Modelle und Erfolgsfaktoren.

Hochschule Luzern HSLU / Schweizerischer Gewerbeverband sgV (2025): Schweizer Unternehmen befürchten Anstieg Cyberkriminalität – Crime-Studie 2025, 31.10.2025. URL: <https://www.hslu.ch/de-ch/hochschule-luzern/ueber-uns/medien/medienmitteilungen/2025/10/31/crime-studie/> (Abruf: 27.06.2026).

IBM Think Topics (2024): Was ist Datenverarbeitung? <https://www.ibm.com/think/topics/data-processing> (Abruf 01.07.2026)

INKUBIT (2025): Zapier vs. Power Automate, Make & N8N – Vergleich der Automatisierungsplattformen.

Inside IT (2024), Statt auf SAP setzt Lidl auf eine Individuallösung. <https://www.inside-it.ch/statt-auf-sap-setzt-lidl-auf-eine-individualloesung-20240503> (27.06.2026)

ISBC (Internation Business Commnincation Standards) (2025). <https://www.ibcs.com/de/resource/horizontal-waterfall-chart/> (Abruf: 27.06.2026)

International Organization for Standardization (2023): ISO/IEC 27035-1:2023 – Information technology – Information security incident management – Part 1: Principles and process. Genf. URL: <https://www.iso.org/standard/78973.html> (Abruf: 27.06.2026).

KMU Admin: Neues Datenschutzgesetz revDSG. <https://www.kmu.admin.ch/kmu/de/home/fakten-trends/digitalisierung/datenschutz/neues-datenschutzgesetz-rev-dsg.html> (Abruf 01.07.2026)

McKinsey: Cloud Economics Study – TCO-Analyse Cloud vs. On-Premise.

National Institute of Standards and Technology NIST (2025): Special Publication 800-61 Revision 3 – Incident Response Recommendations and Considerations for Cybersecurity Risk Management. URL: <https://csrc.nist.gov/pubs/sp/800/61/r3/final> (Abruf 27.06.2026).

Reis, Joe; Housley, Matt (2022): Fundamentals of Data Engineering. Plan and Build Robust Data Systems. Sebastopol: O'Reilly.

SAP (2024): Was ist Robotic Process Automation (RPA)? <https://www.sap.com/resources/what-is-rpa> (Abruf 01.07.2026)

SAP(2025): Was ist ERP? unter <https://www.sap.com/germany/products/erp/what-is-erp.html> (Abruf 27.06.2026)

Schweizerischer Bundesrat (2002): Verordnung über die Führung und Aufbewahrung der Geschäftsbücher (GeBüV) vom 24. April 2002, SR 221.431, Art. 3. URL: <https://www.fedlex.admin.ch/eli/cc/2002/216/de> (Abruf: 01.07.2026).

Schweizerischer Bundesrat (2022): Verordnung zum Bundesgesetz über den Datenschutz (Datenschutzverordnung, DSV) vom 31. August 2022, SR 235.11, Anhang 1. URL: <https://www.fedlex.admin.ch/eli/cc/2022/568/de> (Abruf: 27.06.2026).

Stoiber, R. (2023): Schweizer Datenschutzgesetz (DSG), 29.08.2023. URL: <https://regina-stoiber.com/2023/08/29/schweizer-datenschutzgesetz-dsg/> (Abruf: 27.06.2026).

swiDOC (o. J.): Digitale Archivierung im Finanzwesen. URL: <https://www.swidoc.ch/de/finanzwesen> (Abruf: 27.06.2026).

Weissenberg Group (2022–2024): Was ist Robotic Process Automation? / Die 10 grössten Herausforderungen bei der RPA-Implementierung.

Zacker, Craig (2025), Microsoft 365 Fundamentals. Pearson Education