

Juli 2026

FACILITATING MAGAZINE

INFORMATIONEN AUS DER WELT DES FACILITATING

Mensch-Maschine- Kommunikation

Warum wir im KI-Zeitalter
mehr Facilitating benötigen

Editorial

Liebe Leserin, lieber Leser,

Benutzt ihr bereits KI-Tools? Oder wehrt ihr euch vehement dagegen an? Für beides gibt es gute Gründe und es ist kein Wunder, dass KI polarisiert. Fakt ist jedoch: Die KI ist da.

In dieser 8. Ausgabe des Facilitating Magazine befassen wir uns mit der Frage, wie wir KI im Facilitating und in der Organisationsentwicklung bewusst einsetzen können. Jennifer Konkol stellt die digitale Plattform Rflect vor, die beim Gestalten von Lern-Journeys hilft. Daniel Boos, Gudela Grote und Lena Schneider machen einen Deep-Dive in Workshop-Methode ALIGN, die die Verantwortung der KI-Nutzung klärt. Adrian Essl beleuchtet die Schattenseiten der KI und fragt, was Kinderschutz im digitalen Zeitalter bedeutet. Barbara Backhaus vergleicht Digital-Marketing-Agenturen mit High Reliability Organisationen und Ronald Greber führt uns in die Kunst des Go-Spielens ein. Kurz: Dieses Heft

ist mal wieder ein bunter Mix aus Reportagen, Essays und Hintergrundartikeln.

Wir sind überzeugt, dass gerade in turbulenten Zeiten, wie wir sie im Moment erleben, eine differenzierte Auseinandersetzung mit Technologie zwingend nötig ist. KI ist ein Werkzeug, das uns Angst machen kann oder das wir nutzen können. Sie kann eine Demokratie ins Schwanken bringen oder in den Händen der richtigen Menschen in Stabilisierungsprozessen helfen.

Wir freuen uns auf eine angeregte Auseinandersetzung: Schreibt uns gern auf LinkedIn oder Instagram – oder schickt uns eine Email: Was haltet ihr von KI? Was gefällt euch? Was macht euch skeptisch? Und: Was haltet ihr von dieser Ausgabe des Magazins?

Viel Freude beim Lesen wünschen euch

Inhaltsverzeichnis

KI, Demokratie und die Illusion der Kontrolle	S. 4
Der Moment nach dem Workshop	S. 6
Wer ist in der Pflicht, wenn die KI irrt?	S. 11
Kluge Facilitating-Strategien	S. 16
High Reliability Organizing	S. 20
KI zielt auf das Innere	S. 23
KI – Was sollen wir tun?	S. 27
Im Dialog	S. 28
KI im Training und Coaching	S. 31
Veranstaltungen	S. 33
Impressum	S. 35



Christina Brändli
Verein Facilitating Schweiz/Bern



Renate Franke
school of facilitating, Berlin

Du hast ein Thema und möchtest für uns schreiben?

Melde dich bei info@kreativeloesungswege.ch

Wir freuen uns über neue Autorinnen und Autoren!

KI, Demokratie und die Illusion der Kontrolle

Warum Organisationen im KI-Zeitalter nicht effizienter, sondern ehrlicher werden sollten.
Barbara Backhaus

Wir diskutieren derzeit intensiv über Künstliche Intelligenz und damit verbundene Bias, Transparenz, Regulierung und Verantwortung. Das ist alles wichtig, aber gleichzeitig greift es zu kurz. Denn das eigentliche Problem liegt nicht in der Technologie, sondern in unserer tiefen organisationalen Sehnsucht nach: Kontrolle, Eindeutigkeit und Vorhersagbarkeit. KI verstärkt genau diese Sehnsucht – und macht sie zugleich unerreichbar. Was dabei entsteht, ist eine paradoxe Situation: Je datengetriebener Organisationen werden, desto weniger verstehen

sie, was sie tun. Und genau hier beginnt die Frage nach Demokratie neu.

Organisationen als psychodynamische Systeme

Organisationen, das wird oft missverstanden, sind keine Maschinen. Unser Organisationsverständnis ist gemeinhin ein zentraler blinder Fleck in der aktuellen KI-Debatte, denn Organisationen werden oft als steuerbare Systeme, rationale Entscheidungsapparate oder optimierbare Prozesse behandelt. Aus einer psychodynamischen Perspektive ist das naiv. Organisationen sind vielmehr:

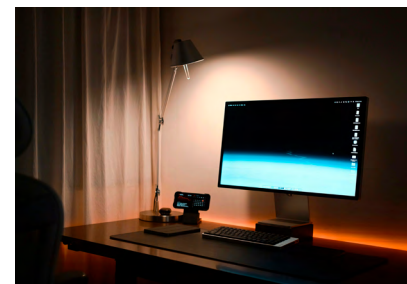
- » Musterstabilisierungsmaschinen
- » Konfliktverarbeitungssysteme
- » Räume kollektiver Unbewusstheit

Was Organisationen tun, folgt nicht nur Logik, sondern auch Angst, Entlastung und Gewohnheit. KI verändert diese Dynamik nicht. Sie verschiebt sie und macht sie dadurch schwerer sichtbar.

Entlastung von Verantwortung

Ein typischer Satz in Organisationen lautet heute: «Die Daten zeigen das.» Das klingt harmlos, ist aber hochpolitisch. Denn dieser Satz tut etwas sehr Spezifisches: Er verschiebt Verantwortung weg von Menschen, hin zu Systemen. Psychodynamisch betrachtet ist das attraktiv, denn Verantwortung

ist anstrengend: Sie erzeugt Unsicherheit und oft auch Konflikte. Und: Sie fordert eine klare Positionierung. Das fällt vielen Menschen schwer. Algorithmen bieten hier eine perfekte Entlastung: Sie treffen Entscheidungen, denen wir folgen können, ohne uns rechtfertigen zu müssen oder uns zu exponieren. Das Problem dabei ist jedoch, dass Organisationen dadurch nicht zuverlässiger, sondern träger werden.



Wenn Systeme entscheiden, wird Verantwortung unsichtbar.

High Reliability beginnt, wo Bequemlichkeit endet

High Reliability Organizing (HRO) beschreibt Organisationen, die unter Unsicherheit zuverlässig handeln. Was dabei oft übersehen wird: Diese Organisationen sind nicht effizient im klassischen Sinn. Sie sind wach. Das heißt, sie trauen einfachen Erklärungen nicht, suchen aktiv nach Fehlern und scheuen keine Ambiguität. Das ist kein strukturelles, sondern ein zutiefst kulturelles Phänomen. Und vor allem: Es ist anstrengend!



Die KI-Hand lenkt. Aber wer lenkt die Hand?

Der deutsche Organisationsberater und Psychotherapeut Klaus Eidenschink hat den Satz geprägt: «Organisationen haben keine Probleme – sie haben Muster.» Diese

Demokratie beginnt nicht bei der Mitbestimmung

Wenn wir über Demokratie im Kontext von KI sprechen, denken wir oft an Beteiligung und Trans-

«Organisationen haben keine Probleme – sie haben Muster.»

(Klaus Eidenschink)

Muster erfüllen Funktionen: Sie reduzieren Komplexität, sichern Anschlussfähigkeit und vermeiden dadurch oft Irritation. Und: Sie sind erstaunlich stabil. Alle, die schon einmal eine Gewohnheit ändern wollten, wissen, wovon ich spreche! KI verstärkt durch Standardisierung, Automatisierung und Wiederholung erfolgreicher Muster diese Stabilität. Das klingt effizient und ist es auch – aber gleichzeitig ist es gefährlich, denn: Was stabil bleibt, wird leicht unsichtbar. Und was unsichtbar ist, kann nicht hinterfragt werden.

Facilitation als Störung

Facilitation wird oft missverstanden als Moderation oder Konsensherstellung. Im Kontext von KI und Demokratie reicht das nicht. Wenn Organisationen zuverlässiger werden sollen, braucht es etwas anderes: Gezielte Störung. Das bedeutet konkret:

- » Fragen, die nicht ins Schema passen
- » Perspektiven, die irritieren
- » Situationen, die Unsicherheit erzeugen

Dabei geht es nicht darum, Chaos zu produzieren, sondern darum, Wahrnehmung überhaupt wieder möglich zu machen. Denn: Ohne Irritation keine Aufmerksamkeit. Und ohne Aufmerksamkeit keine Verantwortung.

parenz. Das ist zwar wichtig, reicht aber oft nicht aus. Demokratie beginnt nämlich schon viel früher, mit der Frage: Wer darf und kann wahrnehmen, was passiert? In vielen Organisationen ist Wahrnehmung durch Hierarchien und Rollen und daran geknüpfte implizite Erwartungen bereits eingeschränkt. KI kann diese Einschränkungen verstärken, indem sie bestimmte Daten privilegiert und damit andere Perspektiven ausblendet. Facilitating hat hier eine zentrale Funktion: Sie erweitert den Raum des Wahrnehmbaren. Genau das ist demokratische Praxis.

Die Zumutung der echten Achtsamkeit

Kollektive Achtsamkeit, das Herzstück von HRO, klingt harmlos, fast schon weich. Doch in Wirklichkeit ist sie eine Zumutung.

Sie bedeutet, Dinge zu sehen, die man lieber nicht sehen würde, und Zweifel auszuhalten, wo Klarheit erwartet wird. In einer KI-dominierten Welt wächst diese Zumutung. Wenn wir Organisationen im KI-Zeitalter ernstnehmen, dann geht es nicht darum, bessere Tools einzuführen, mehr Daten zu sammeln und zu erwarten, dadurch schnellere Entscheidungen zu treffen. Vielmehr geht es darum:

- » Wahrnehmung zu organisieren: Wer sieht was – und wer nicht?
- » Muster sichtbar zu machen: Was wiederholt sich – und warum?
- » Irritation zu ermöglichen: Wo können wir uns selbst unterbrechen?
- » Verantwortung zurückzuholen: Wer steht für Entscheidungen ein?

Weniger KI, mehr Mut

KI wird bleiben. Und sie wird besser werden. Die entscheidende Frage ist also nicht, ob wir sie nutzen, sondern wie. Wenn Organisationen versuchen, Unsicherheit durch Technologie zu eliminieren, werden sie scheitern. Wenn sie beginnen, Unsicherheit gemeinsam wahrzunehmen und zu verarbeiten, entsteht etwas anderes: Zuverlässigkeit. Und vielleicht sogar eine neue Form organisationaler Demokratie.



Wer nichts aufwühlt, verändert nichts.

Der Moment nach dem Workshop

Wie KI den Transfer in Bewegung hält

Gute Workshops erzeugen Energie, aber Wirkung entsteht erst danach. Wie können KI-gestützte Plattformen den Transfer in den Alltag unterstützen, Reflexion verstetigen und Entwicklung sichtbar machen? Und wo liegen die Grenzen?

Jennifer Konkol

Wie schön ist es, wenn wir am Ende von Workshops feststellen, dass die Energie gut ist! Dass der Kopf verstanden und das Herz gefühlt hat! Wie schön, wenn die Teilnehmenden rausgehen und sagen: «Das war gut, ich nehme etwas mit für mich.»

Und dann kommt der Alltag und damit die Phase, in der sich Wirkung zeigt – oder eben nicht.

Die Sache mit dem Transfer

Eine Studie der Forscherin Nele Graf aus dem Jahr 2016 mit über 10000 befragten Mitarbeitenden zeigt ein bekanntes Muster: Fast alle verstehen, wie wichtig Lernen ist, aber die Umsetzung im Alltag gelingt nur einem kleinen Teil wirklich konsequent. Selbststeuerung, Dranbleiben und Transfer sind für viele die eigentlichen Herausforderungen. Auch der Psychologe Axel Koch beschreibt in seinem Ansatz der Transferstärke, dass nachhaltige Verhaltensänderung weniger am Wissen scheitert als an Faktoren wie Umsetzungskompetenz, Selbststeuerung und unterstützenden Rahmenbedingungen im Alltag.

Besonders deutlich wird es in Trainings, in denen es nicht nur um

neue Methoden und neues Wissen geht, sondern um ein anderes Arbeiten oder Führen. Der Entwicklungspsychologe Robert Kegan unterscheidet dabei zwischen horizontaler und vertikaler Entwicklung. Die horizontale Entwicklung füllt den Werkzeugkasten mit Kenntnissen, Tools und Methoden, während die vertikale Entwicklung die Person, die den Werkzeugkasten benutzt, mit Haltung sowie inneren Wahrnehmungs- und Denkfähigkeiten verändert.

Und genau das macht es so anspruchsvoll, denn vertikale Entwicklung passiert nicht in einem Workshop. Sie entsteht nachgelagert im Alltag, wenn ich Situationen anders wahrnehme und mich bewusst mit ihnen auseinandersetze. Wenn wir also an vertikaler Entwicklung arbeiten wollen, reicht ein gutes Training nicht. Es braucht Momente im Alltag, in denen ich wieder hinschaue, reflektiere und ausprobieren. Wir als Facilitator:innen versuchen, genau hier anzusetzen.



Durch erlebnisorientiertes Trainingsdesign, durch Peer-Unterstützung, durch Lernbuddys, durch mehrteilige Lernreisen, durch gezielte Reflexions- und Austauschformate im Nachgang. Und trotzdem zeigt sich immer wieder: Transfer bleibt schwierig.

Psychologische Sicherheit und soziale Führungskompetenz in einem stark technisch geprägten Bereich gezielt stärkt. Die Herausforderung dabei war, Themen wie Emotionen, Beziehungsgestaltung und offenes Ansprechen von Fehlern so zu vermitteln, dass sie im All-

Vertikale Entwicklung passiert nicht in einem Workshop. Sie entsteht nachgelagert im Alltag.

Die Sichtbarkeit der Entwicklung ist schwer greif- und messbar. Für mich war diese Erkenntnis der Ausgangspunkt, um neu hinzuschauen und zu prüfen, welche Rolle KI dabei spielen kann.

Rflect als Plattform für innere Entwicklung

Ich kannte Niels Rot, einen der Co-Gründer von Rflect, schon etwas länger. Mich hat immer beeindruckt, wie stark ihn und das Team die Frage antreibt, wie man Inner Development wirklich in den Alltag bringt. Als wir uns auf der Inner Development Goals Conference in Stockholm wiedertrafen, kamen wir ins Gespräch. Dabei wurde die Idee konkret, Rflect für eines meiner Kundenprojekte anzuwenden. Normalerweise wird Rflect vor allem im Hochschulkontext angewendet, aber mit unserem Gespräch wurde der erste Pilot für den Corporate Kontext geboren.

Psychologische Sicherheit im technischen Umfeld verankern

Für ein grosses Schweizer Unternehmen wurde ich als externe Facilitatorin engagiert, um ein Programm zu entwickeln, das psy-

tag des technischen Produktionsumfelds wirklich anschlussfähig sein würden.

Dafür setzten wir ein Pilotprojekt auf, in dem eine Gruppe von 13 Führungskräften ein kompaktes Programm durchläuft. Dieses bestand aus einem Kick-off, einem zweitägigen Training sowie (einige Zeit später) einem weiteren Workshop zur Reflexion und Einordnung. Ergänzend entwickelten wir eine Toolbox, die im Führungsalltag genutzt werden kann. Parallel dazu läuft aktuell eine Evaluation, um die Entwicklung und Wirkung sichtbar zu machen und am Ende entscheiden zu können, ob ein Roll-out sinnvoll ist.

Das Programm wird intern stark durch die Fachbereiche BGM, Organisationsentwicklung und Arbeitssicherheit getragen und von zwei externen Facilitatorinnen (Lisa Ritter und mir) begleitet. Zentral ist auch das Commitment der obersten Führungsebene des Bereichs, die das Vorhaben aktiv mitträgt. Zwei Sponsor:innen aus ihren Reihen nehmen ausserdem selbst daran teil.

Was ist Rflect?

Eine digitale Plattform, mit der sich Lern- und Entwicklungsprozesse als strukturierte Lern-Journeys gestalten und begleiten lassen. Ihr Fokus liegt auf persönlichem Wachstum und Inner Development. Die Plattform kombiniert Reflexion mit kleinen Assessments, Lernziel-Formulierung und Peer-Coaching-Formate, sodass Lernen nicht isoliert, sondern im Austausch und entlang konkreter Entwicklungsziele stattfindet.

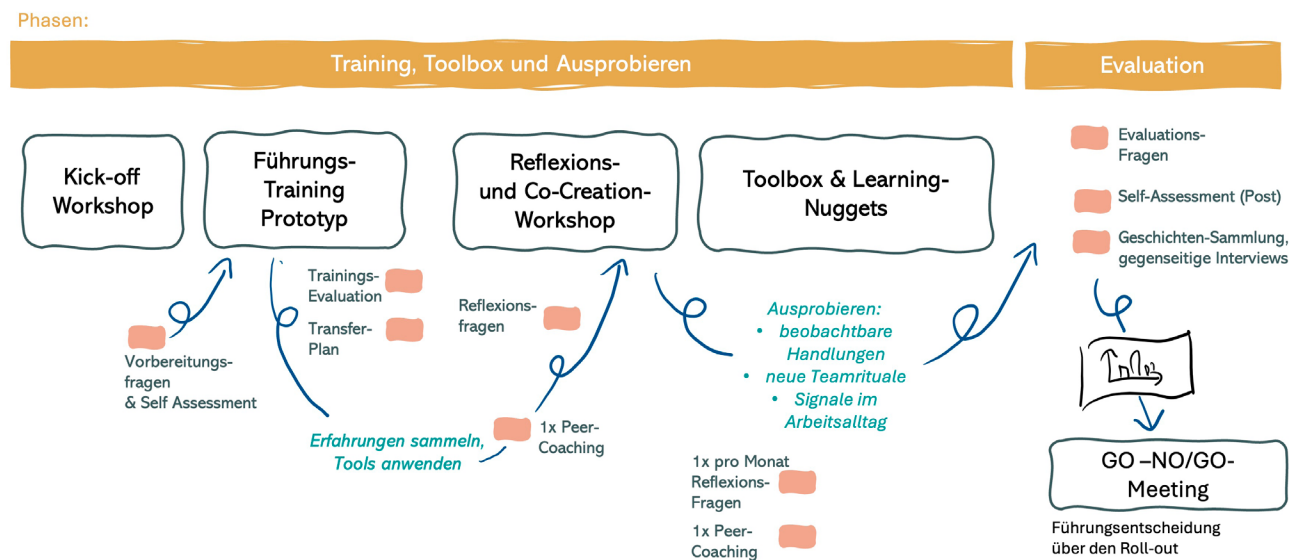
Eine integrierte KI unterstützt sowohl Lernende als auch Trainer:innen: Sie hilft, wirksame Lernziele zu formulieren, Muster in Reflexionen und Daten zu erkennen, Fortschritte sichtbar zu machen und gezielt Hinweise zu geben, wo Vertiefung oder Unterstützung sinnvoll ist. Bisher wird Rflect an mehr als 35 europäischen Hochschulen für die Unterstützung von Inner Development eingesetzt.

rflect.ch

Wie Rflect den Transfer konkret unterstützt

Zu Beginn unterstützt Rflect bei der Vorbereitung. Teilnehmende setzen sich mit ersten Fragen auseinander und führen ein Self Assessment durch. Dieses Self Assessment ist an die Kompetenzen gekoppelt, die unser Auftraggeber bei seinen Mitarbeiter:innen entwickeln möchte. Dadurch entsteht schon vor dem Training eine erste Reflexion.

Direkt nach dem Training wird die Plattform genutzt, um Feedback



Die Abbildung zeigt die Roadmap des Piloten. Die rot markierten Punkte machen sichtbar, an welchen Stellen Rflect als verbindendes Element zwischen den einzelnen Formaten eingesetzt wird.

zu geben und die eigene Transferfähigkeit einzuschätzen, angelehnt an das Konzept von Axel Koch. Ein Teil davon ist ein klassischer Fragebogen. Dazu kommen offene Reflexionen über das Gelernte, die von den Teilnehmenden eingetippt oder ausgesprochen werden können. Mit Hilfe einer KI werden die Reflexionen analysiert. Dadurch handelt es sich beim Feedback weniger um ein «Add-on»,

Gleichzeitig setzen sich die Teilnehmenden eigene Lernziele. Rflect unterstützt sie mittels KI dabei, diese Ziele zu schärfen und auf eine Flughöhe zu bringen, auf der sie wirklich umsetzbar und gleichzeitig motivierend sind. Das spart im Training viel Zeit und führt dazu, dass die Ziele konkreter werden. Zusätzlich wird der Transferplan in konkrete Schritte heruntergebrochen: Welche kleinen Veränderungen

Im weiteren Verlauf unterstützt die Plattform vor allem das Dranbleiben. Ein- bis zweimal pro Monat gibt es Reflexionsimpulse. Diese können individuell bleiben oder – je nach Einstellung – mit Facilitator:innen oder in der Gruppe geteilt werden. Das schafft Orientierung, aber auch gegenseitige Inspiration.

Integrierter Workflow statt klassischer Chatbot

Das Peer-Coaching wird ebenfalls über die Plattform initiiert. Teilnehmende werden gezielt miteinander verbunden. Dabei kann beispielsweise darauf geachtet werden, dass eher aktive Teilnehmende mit weniger aktiven Personen in Austausch gebracht werden. Ich kann die Gruppen aber auch so einteilen, dass sie sowohl team- als auch hierarchieübergreifend gemischt sind. Die Teilnehmenden erhalten direkt über die Plattform eine einfache Struktur für das Gespräch. Im Anschluss reflektieren sie das Gespräch wieder in Rflect und halten fest, was sie daraus mitnehmen.

Die KI ist kein klassischer Chatbot, sondern ein integrierter Workflow, der die einzelnen Schritte der Lernreise organisch miteinander verbindet.

sondern vielmehr um eine direkte Ableitung aus einer persönlich-relevanten Reflexion. Sie wird dadurch viel aussagekräftiger.

gen sind realistisch? Was könnten erste «Ein-Prozent-Schritte» sein? Welche Hürden könnten auftauchen? Und: Wie gehe ich damit um?

Über den ganzen Prozess haben die Teilnehmenden die Möglichkeit, konkrete Mini-Actions zu definieren, also kleine, realistische nächste Schritte für ihren Alltag. Auch hier unterstützt die KI dabei, diese Schritte klar und handhabbar zu formulieren, und erinnert an die Umsetzung. Die KI ist also kein klassischer Chatbot wie ChatGPT oder Claude, sondern ein integrierter Workflow, der die einzelnen Schritte der Lernreise organisch miteinander verbindet. Begleitend dazu haben wir es so eingestellt, dass die Nutzung der Toolbox in der Lernreise immer wieder aufgegriffen wird und auch die Vor- und Nachbereitung der weiteren Workshops über die Plattform stattfindet.

Der KI Teaching Assistant im Hintergrund

In der Abschlussphase wird Rflect für die Evaluation genutzt. Dazu gehören:

- » ein erneutes Self-Assessment
- » klassische Evaluationsfragen
- » qualitative Einblicke, z.B. über gegenseitige Interviews und gesammelte Geschichten aus dem Transfer

Diese verschiedenen Perspektiven fließen zusammen und bilden die Grundlage für die Entscheidung, ob und wie das Programm skaliert wird.

Für uns als Facilitator:innen spielt der KI Teaching Assistant im Hintergrund eine wichtige Rolle. Er hilft dabei, Muster in Reflexionen und Rückmeldungen zu erkennen und besser zu verstehen, wo die Teilnehmenden aktuell stehen: Was ist schon angekommen? Wo gibt es noch Unklarheiten? Wo braucht es vielleicht einen anderen Impuls?

Darüber hinaus unterstützt der KI Teaching Assistant die Teilnehmenden

den dabei, die Learning Journey aufzusetzen und weiterzuentwickeln. Er gibt Hinweise, welche nächsten Schritte sinnvoll sein

» **Bessere Einblicke in die Gruppe:** Als Facilitator:innen sehen wir auch im Transferprozess schnell, wo die Gruppe

Mit Rflect passiert Wirkungsmessung eher als Nebenprodukt von etwas, das echten Mehrwert für die Teilnehmenden bringt.

könnten, wo Vertiefung hilfreich wäre oder wie einzelne Lernimpulse noch besser gestaltet werden können.

Ein Zwischenfazit

Stand Sommer 2026 sind wir mitten in der Transferphase des Projekts. Entsprechend ist das hier ein Zwischenstand.

Was gut funktioniert:

» **Dranbleiben im Alltag:** Die regelmässigen Impulse und Erinnerungen sorgen dafür, dass Themen nicht verschwinden. Nur eine von 13 Personen nutzt Rflect bisher gar nicht, über die Hälfte hat mehr als 80 % der Reflexionen erledigt. Die Beiträge sind inhaltlich tief: Teilnehmende berichten von konkreten Schritten und ersten Erfolgen, die sie bereits umgesetzt haben.

» **Technische Einführung:** Die Einführung im Corporate Kontext lief erstaunlich reibungslos. Es gab keine Probleme mit Firewalls, wenig Erklärungsbedarf, die Plattform wurde intuitiv verstanden.

steht: Wer ist mehr oder weniger aktiv? Wo hakt es? Der KI Teaching Assistant hilft zusätzlich, Muster in Reflexionen zu erkennen und einzuordnen, was schon verstanden wurde und wo es noch Klärung braucht.

» **Sichtbarkeit schafft Verbindung:** Wenn Beiträge für alle sichtbar geteilt werden, kann man voneinander lernen und es wird sichtbar, dass andere auch aktiv sind. Das schafft Verbindung und wirkt motivierend. Die oberste Führungsebene kann hier aktiv als Vorbild agieren und damit zeigen: Wir machen mit!

» **Weniger Organisationsaufwand:** Besonders beim Peer-Coaching und bei der Evaluation zeigt sich ein klarer Vorteil. Das Aufsetzen passiert schnell und strukturiert, mit deutlich weniger Aufwand als sonst.

» **Lernerfolg sichtbar machen:** Ein grosser Mehrwert liegt in den Auswertungen: Entwicklung wird über die Zeit sichtbar, nicht nur punktuell. Das schafft Klarheit für Organisationen, wirkt motivierend für Teilnehmende

und hilft Facilitator:innen, die Lernreise gezielt anzupassen.

Das Team von Rflect arbeitet derzeit ausserdem daran, Entwicklung über die Analyse der Veränderungen in der genutzten Sprache sichtbar zu machen. Das ist ein spannendes Potential, denn Sprache ist eng mit unserem Denken verbunden: Wie wir etwas formulieren, zeigt oft, wie wir es innerlich strukturieren. Sprache als Spiegel für sich verändernde Denkmuster zu nutzen und daraus Hinweise auf vertikale Entwicklung abzuleiten, scheint ein vielversprechender Weg, Wirkung auf dieser Ebene sichtbar zu machen.

Wo es herausfordernd bleibt:

- » **Ein weiteres Tool im Alltag:** Neben bestehenden Systemen wie Teams oder SharePoint ist es eine Hürde, ein weiteres Tool aktiv zu nutzen. Hier braucht es bewusste Einbettung. Hilfreich wäre es, wenn man die Trainingsunterlagen, Tools und Infos auch bei Rflect ablegen könnte, sodass die Teilnehmenden alles an einem Ort finden.
- » **Die richtige Taktung finden:** Wie viele Impulse sind sinnvoll? Wann passen sie gut? Wie lang dürfen sie sein?

» **Nutzen vs. Aufwand:** Hilft die Plattform auch aus Sicht der Teilnehmenden im Alltag spürbar oder wird sie als zusätzliche Aufgabe erlebt? Das wird sich erst in der abschliessenden Evaluation zeigen.

In einer Zeit, in der Organisationen sehr genau hinschauen, wo sie investieren, wird die Frage nach Wirkung immer wichtiger. Besonders bei vertikaler Entwicklung ist das herausfordernd, weil sie sich nicht so leicht messen und greifen lässt.

Bisher habe ich das oft als Spannungsfeld erlebt: der Anspruch, Wirkung sichtbar zu machen, und gleichzeitig die Gefahr, dass Messung Druck erzeugt oder zum Selbstzweck wird. Mit Rflect passiert Wirkungsmessung eher als Nebenprodukt von etwas, das echten Mehrwert für die Teilnehmenden bringt. Das macht den Umgang damit deutlich leichter und gibt dem Thema Messung eine andere Qualität – weniger kontrollierend, mehr unterstützend. Und genau das finde ich aus Facilitator:innen-Perspektive besonders stark.

Quellen

Graf, N.; Gramß, D. und Heister, M. (2016) LEKAF-Studie *Gebrauchsanweisung fürs lebenslange Lernen*, Düsseldorf: Vodafone Stiftung.
Kegan, R. (1982). *The evolving self: Problem and process in human development*. Harvard University Press.
Koch, A. (2023). Die Transferstärke-Methode: Befähigung zum selbstverantwortlichen Lerner. In *Lernen im Zeitalter der Digitalisierung: Einblicke und Handlungsempfehlungen für die neue Arbeitswelt* (pp. 77-95). Wiesbaden: Springer Fachmedien Wiesbaden.

Jennifer Konkol

begleitet seit 2008 Teams, Führungskräfte und Organisationen in Entwicklungs- und Transformationsprozessen. Ihre Vision: Unternehmen dabei zu unterstützen, ihre inneren Funken, d.h. ihre Potenziale und Stärken, lebendig werden zu lassen.

funkenspruehen.ch

Wer ist in der Pflicht, wenn die KI irrt?

Verantwortung und Kontrolle in KI-Projekten

Die Workshop-Methode ALIGN ist ein Ansatz zur partizipativen Klärung von Verantwortung und Kontrolle bei der Entwicklung und Nutzung von KI. «Human in the loop» klingt nach Kontrolle. Aber was bedeutet das wirklich? Und: Reicht es?

Daniel Boos, Gudela Grote & Lena Schneider

Ein Krankenhaus setzt KI-gestützte Software zur Erkennung von Lungenembolien ein. Das Personal hat eine Basisschulung erhalten, kennt aber die Funktionsweisen und Einschränkungen des Systems nicht. Als ein Patient mit Atemnot in die Notaufnahme kommt, erkennt die KI die Lungenembolie nicht. Die seit vielen Stunden im Dienst stehende Radiologin folgt der Einschätzung des Systems und der Patient wird entlassen. Am nächsten Tag stirbt der Patient.

Wer trägt nun die Verantwortung?

Die Radiologin? Sie hat die Verantwortung für die Analyse, kennt aber die zugrundeliegenden Prozesse und möglichen Fehler oder Unsicherheiten nicht und hat daher keine Möglichkeit, die Darstellung kritisch einzuschätzen.

Der Hersteller? Er hatte Kontrolle über das Design der Anwendung, nicht aber über deren Einsatz in der Notaufnahme.

Das Spital? Es hat den Einsatz der Technologie proaktiv gefördert und über die Einführung entschieden.

Die Zulassungsbehörde? Sie hat basierend auf den ihr zur Verfügung gestellten Informationen über die Zulassung und den Einsatz in Spitälern entschieden.

Vor allem: Hätten diese Fragen nicht schon vor dem Einsatz geklärt sein müssen?

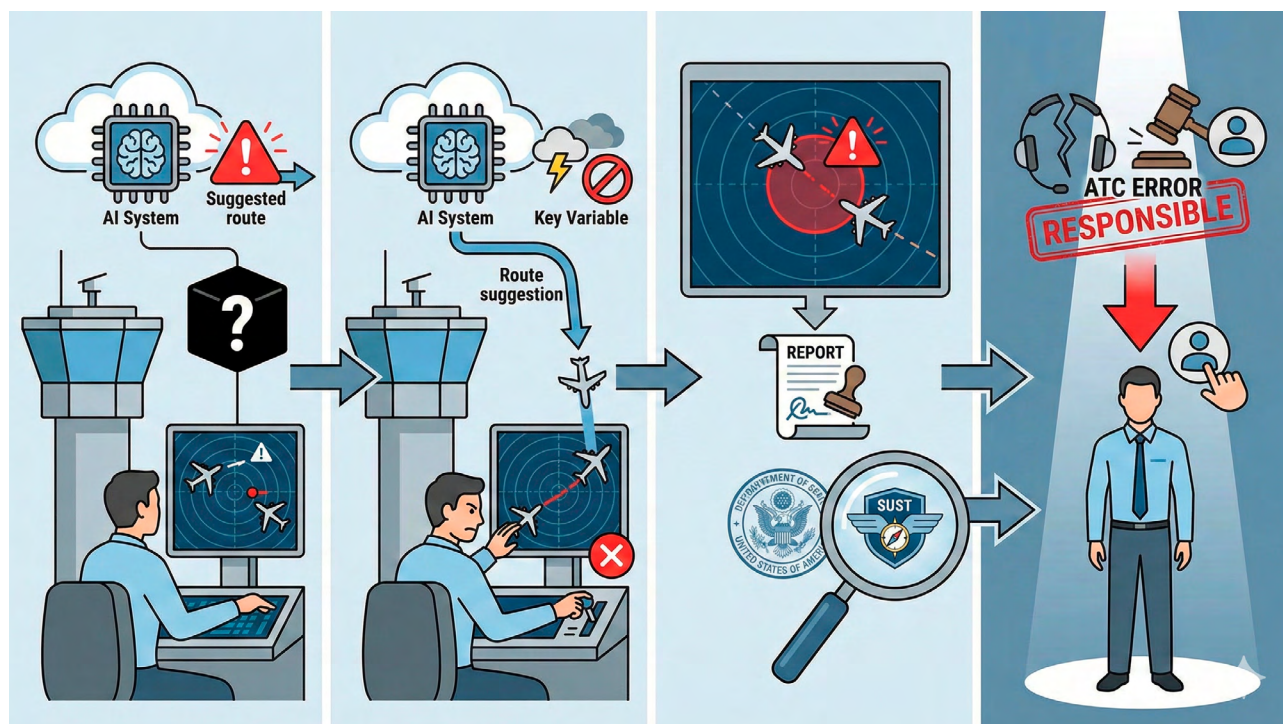
Dieses Szenario ist fiktiv, könnte sich aber genau so abspielen. Es stammt von ALIGN, einem partizipativen Gestaltungsansatz zur Passung von Verantwortung und Kontrolle in der Entwicklung und Nutzung von KI. Wie können Fragen zu Verantwortung und Kontrolle geklärt werden, bevor etwas schief läuft? Was braucht es, damit Verantwortung und Kontrolle für alle Akteur:innen zusammenpassen?

Passung von Verantwortung und Kontrolle

Bei der Auswahl, Entwicklung und



Eine erschöpfte Radiologin folgt mitten in der Nacht der KI-Einschätzung – ein formales «Human in the loop», dem die nötigen Voraussetzungen fehlen.



Visualisiertes Beispielszenario aus dem Bereich «Luftfahrt»

Einführung von KI in Unternehmen braucht eine klare Passung zwischen Verantwortung und Kontrolle für alle Beteiligten. Kontrolle ist die Fähigkeit von Akteur:innen, Arbeitsprozesse und deren Ergebnisse zu beeinflussen, um gewünschte Ziele zu erreichen oder unerwünschte Konsequenzen zu vermeiden. Dafür sind mehrere Elemente nötig: die Fähigkeit, auf Situationen Einfluss zu nehmen, Transparenz über den Systemstatus, Nachvollziehbarkeit von Abläufen und die Vorhersehbarkeit von Ergebnissen.

Die Verantwortung aller Akteur:innen spielt dabei eine wesentliche Rolle. Sie bezeichnet die Pflicht der verantwortlichen Person, ihre Handlungen gegenüber anderen sichtbar zu machen, sie zu begründen und Zuständigkeiten bis zur Haftbarkeit wahrzunehmen. Dabei gilt der Gestaltungsgrundsatz, dass Menschen nur für Handlungen und deren Folgen verantwortlich gemacht werden können, über

die sie effektiv Kontrolle haben. Die Passung von Verantwortung und Kontrolle rückt also ins Zentrum und die beiden Konzepte bedingen sich gegenseitig.

Mehrere Eigenschaften von KI beeinflussen die Passung, weshalb diese proaktiv und iterativ adressiert werden muss. KI-Systeme sind probabilistisch. Das heisst, sie erzeugen nicht immer dasselbe Resultat. Wir kennen das beispielsweise von KI-Chats: Das dahinterstehende Modell kann auf die gleiche Frage unterschiedliche Antworten liefern. Das ist herausfordernd, wenn es um risikobehaftete Entscheide geht. Für Nutzer:innen und abgeschwächt auch für KI-Entwickler:innen gleicht die Funktionsweise von KI einer Black Box und ist somit wenig durchschaubar. KI-Systeme können ausserdem selbstständig lernen und verändern deshalb ihre Fähigkeiten und Antworten. Die Systeme sind zudem hochgradig vernetzt, viele technische und menschliche

Akteur:innen sind an Entwicklung, Zulassung, Implementierung und Nutzung beteiligt.

Das Trugbild der menschlichen Aufsicht

Zunehmend agieren KI-Systeme auch als eigenständige Agenten, wodurch die Frage nach Verantwortung noch komplizierter wird. Als Lösung wird vorschnell vom «Human in the loop» gesprochen, wie es etwa der EU AI Act und viele weitere Regelwerke fordern. Oft wird dieses Konzept weder im Detail erklärt, noch wird ausgeführt, wie es konkret erreicht werden könnte.

Dabei müssen wir zwischen einem eher formalen und einem wirklichen «Human in the loop» unterscheiden. Ein formaler Einbezug menschlicher Akteur:innen bedeutet zum Beispiel, dass sie nur das Ergebnis bestätigen oder ablehnen können. Implizit wird dadurch die Verantwortung zurück an die Menschen gegeben, obwohl

viele Entscheidungsschritte durch das System schon vorgängig determiniert wurden. Es kommt zu einer Weiterführung der Ironie der Automatisierung: Menschen überwachen von aussen, werden bei Ausnahmen involviert – stehen dann aber häufig unter Zeitdruck, sind zu schlecht trainiert oder haben zu wenig Kontrolle über das System, um ausreichend Einfluss nehmen zu können. Ein wirkliches «Human in the loop» heisst dagegen, dass Menschen Zeit, Fachwissen, Verständnis des Systems und tatsächliche Eingriffsmöglichkeiten haben.

Die fehlende Passung ist also kein rein technisch zu lösendes Problem, sondern betrifft die Gestaltung der Rollen, der Aufgaben und der Organisation der Arbeitsprozesse. Die fehlende Passung entsteht nicht erst durch den Einsatz des KI-Systems, sondern schon während der Entwicklung. Wichtige Entscheide werden in der Entwicklungsphase vorgespart, etwa

durch die Wahl des umzusetzenden Use Case, der dazu benötigten KI-Technologie oder der genutzten Trainingsdaten. Regulatorische Vorgaben definieren zusätzliche Anforderungen bezüglich Verantwortlichkeit. Fehlende Passung betrifft zudem nicht nur die Nutzer:innen der Technologie und deren Vorgesetzte, sondern ein zunehmend wachsendes Netzwerk von involvierten Akteur:innen sowohl in der Entwicklung des Systems (z.B. Modellentwicklung, Aufbereitung und Validierung von Trainingsdaten, Modelltests) als auch in dessen Lizenzierung und Marktüberwachung.

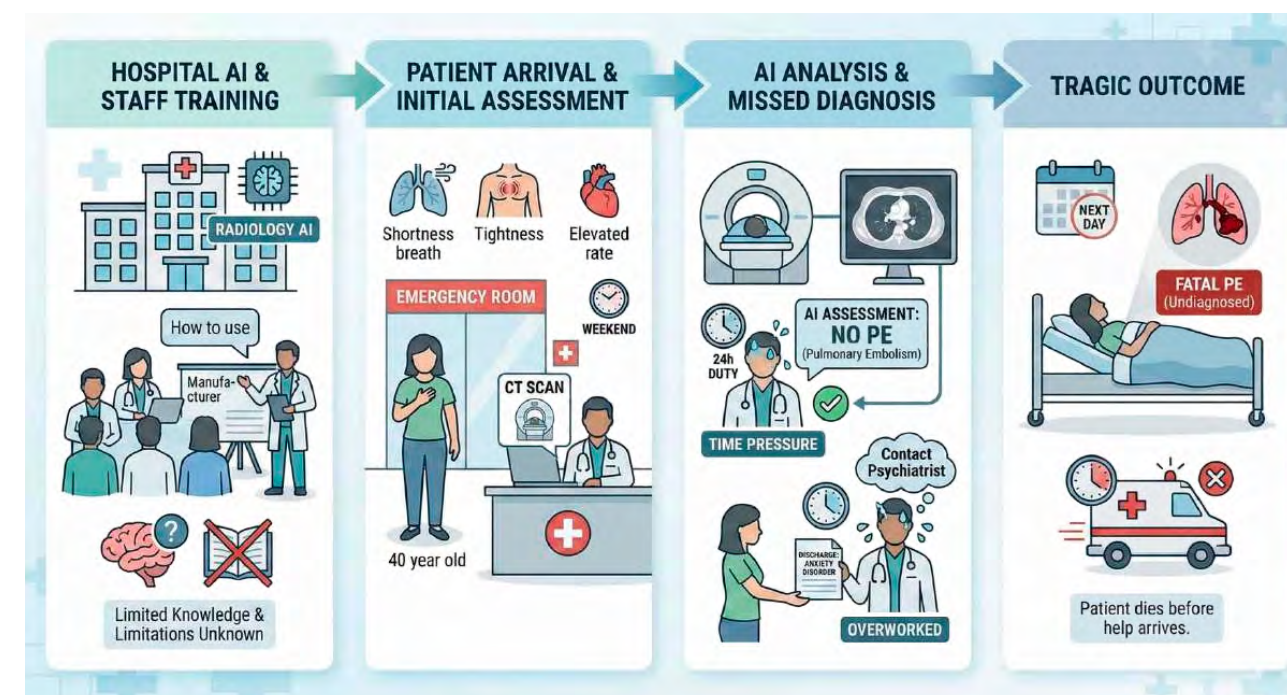
ALIGN als Lösungsansatz

Da bei KI-Vorhaben eine Vielzahl von heterogenen Anspruchsgruppen beteiligt ist, die unterschiedliches Wissen und unterschiedliche Interessen und Einflussmöglichkeiten haben, werden Ansätze benötigt, die eine partizipative Gestaltung ermöglichen. Ziel ist es, die verschiedenen Perspekti-

ven zusammenzubringen und proaktiv Massnahmen zur besseren Passung von Verantwortung und Kontrolle für alle beteiligten Akteur:innen zu gestalten, bevor etwas schief läuft.

ALIGN ist eine Workshop-Methode, die an der Professur für Arbeits- und Organisationspsychologie im Rahmen der ETH Mobilitätsinitiative entwickelt und mit der SBB erprobt wurde. Zur Zielgruppe gehören vier bis sechs Personen. Sie repräsentieren die verschiedenen Anspruchsgruppen, die im geplanten KI-Vorhaben einer Organisation involviert oder davon betroffen sind. Ziel des Workshops ist es, für alle Akteur:innen die vorgesehenen Verantwortungen und Kontrollmöglichkeiten sowie allfällige Lücken in deren Passung sichtbar zu machen und gemeinsam Lösungen zu erarbeiten. Der Workshop benötigt etwa zwei Stunden, plus Vor- und Nachbereitungszeit.

Der Workshop ist in drei Phasen



	Muster 1 Tiefe Komplexität des KI Systems Volle Kontrolle und Verantwortung für KI-Nutzende	Muster 2 Mittlere Komplexität des KI-Systems Teilweise Kontrolle für KI-Nutzende kombiniert mit voller Kontrolle und Verantwortung für KI-Entwickelnde	Muster 3 Hohe Komplexität des KI Systems Teilweise Kontrolle und volle Verantwortung für KI-Entwickelnde
Kontrolle über KI-Ergebnisse	KI-Nutzende	KI-Nutzende (teilweise)	KI-Nutzende (sehr eingeschränkt)
Verantwortung über KI-Ergebnisse	KI-Nutzendea	KI-Entwickelnde	KI-Entwickelnde
Kontrolle über KI-Funktionsweise	KI-Entwickelnde	KI-Entwickelnde	KI-Entwickelnde (teilweise)
Verantwortung für KI-Funktionsweise	KI-Entwickelnde	KI-Entwickelnde	KI-Entwickelnde Senior Management
Unterstützende organisationale Massnahmen	Handlungsspielraum der KI-Nutzenden stärken	Kontrollmissbrauch durch Nutzende verhindern	Unterstützung von Monitoring- und Feedbacksystemen für kontinuierliches Lernen
KI-Systemkomplexität	Niedrig	Mittel	Hoch

Drei Muster für Passung von Verantwortung und Kontrolle für KI-Vorhaben mit unterschiedlich hoher Komplexität

aufgebaut. Zuerst erhalten die Teilnehmenden eine kurze Einführung in die Konzepte von Verantwortung und Kontrolle sowie die KI-spezifischen Herausforderungen bei deren Passung.

Im zweiten Teil werden Szenarien aus anderen Branchen diskutiert (u.a. Luftfahrt, Bahn, Medizin; mit Experten aus diesen Bereichen validiert), anhand derer gezeigt wird, was passiert, wenn Verantwortung und Kontrolle auseinanderfallen. Es sollen bewusst Szenarien aus anderen Branchen gewählt werden, damit die Teilnehmenden das Thema unvoreingenommen diskutieren können und nicht durch mögliche eigene Interessen beeinflusst werden. Gemeinsam werden dann Gestaltungsideen für diese Szenarien entwickelt. Im dritten Teil wechselt der Fokus auf das eigene KI-Vorhaben: Wer ist beteiligt? In welcher Rolle, mit welcher Verantwortung und mit welchen Kontrollmöglichkeiten? Wo gibt es

Lücken zwischen Verantwortung und Kontrolle? Diese Lücken werden sichtbar gemacht und priorisiert. Die Ideen aus dem zweiten Teil werden genutzt, um daraus konkrete Massnahmen abzuleiten.

Zusätzlich kann die Tabelle mit drei Mustern zur Passung von Verantwortung und Kontrolle als Unterstützung beigezogen werden, da je nach Komplexität die Kontrollmöglichkeiten unterschiedlich sind. So ist bei tiefer Komplexität des KI-Systems der Handlungsspielraum der Nutzenden grösser. Bei hoher Komplexität wiederum ist der Gestaltungsspielraum beim Management. In allen Mustern tragen KI-Entwickelnde oder Lösungsanbieternde Verantwortung.

Partizipativ statt reaktiv
Der Erfolg der Nutzung von ALIGN hängt mitunter von der Facilitation ab. Facilitator:innen können als neutrale Personen darauf achten, dass nicht eigene Interessen

einzelner Anspruchsgruppen dominieren und auch heikle Fragen zu Verantwortung und Kontrolle adressiert werden. Wichtig ist dabei nicht nur der Blick auf Situationen, in denen Nutzende die Verantwortung tragen, ohne über passende Kontrollmöglichkeiten zu verfügen: Relevant sind auch Fehlpassungen, wenn andere Akteur:innen, etwa Entwickelnde oder Lösungsanbieternde, ein System gestalten und damit Kontrolle ausüben, aber keine Verantwortung für die Nutzung tragen. Gerade in vernetzten Umgebungen mit vielen heterogenen Akteur:innen ist partizipative Gestaltung zentral, damit die Lösung von allen getragen werden kann. Beteiligung ist der Schlüssel dazu. Facilitation ist dann erfolgreich, wenn sie die unterschiedlichen Perspektiven, Aufgaben und Interessen produktiv zusammenbringt. Die Methode sollte nicht nur einmal eingesetzt werden, sondern kann mehrfach vor wichtigen Entscheiden bei der

Entwicklung, Einführung oder Weiterentwicklung genutzt werden.

Wer ALIGN nicht in einer vollständigen Version einsetzen kann, hat die Möglichkeit, das Thema auch mit einer kleineren Intervention bei einem Austausch von Stakeholdern aufzugreifen und folgende Fragen zu stellen: Wer ist verantwortlich, wenn das System einen Fehler macht? Wer kann tatsächlich eingreifen? Gibt es Beteiligte, die Verantwortung tragen, aber keine Kontrollmöglichkeiten haben? Gibt es Akteur:innen, die Kontrolle ausüben, aber keine Verantwortung tragen?

KI verändert und automatisiert nicht nur Aufgaben, sondern verdeutlicht auch, wer für die Ergebnisse verantwortlich ist und wer welche Möglichkeiten der Kontrolle hat. Die Frage lässt sich nicht allein technisch lösen. Sie erfordert eine bewusste Gestaltung der Aufgaben, der Kontrollmöglichkeiten und der Verantwortung. Dies geschieht iterativ – bevor Probleme auftreten oder sogar Schaden entsteht, nicht erst danach. Dies gilt besonders in einer Arbeitswelt mit zunehmend verteilt agierenden Menschen und autonomen KI-Agenten.

Quellen
Grote, G., Parker, S. & Crowston, K. (2024). Taming Artificial Intelligence: A Theory of Control–Accountability Alignment among AI Developers and Users. Academy of Management Review
Schneider, L., Boos, D., Theurillat, P. & Grote, G. (2025). A Toolbox for Assessing and Negotiating Accountabilities among Heterogeneous Stakeholders in the Context of AI-based Systems in the Railways. Konferenzbeitrag, Human Factors in Railway (HFRAIL), London, September 2025.
ETH Mobilitätsinitiative, Projekt ExplainAI: csfm.ethz.ch/en/research/projects/explain-ai.html



Template für Fit Assessment

ALIGN

wurde im Rahmen der ETH Mobilitätsinitiative vom Lehrstuhl für Arbeits- und Organisationspsychologie der ETH gemeinsam mit der SBB und Siemens entwickelt und erprobt.

Die Autor:innen

Dr. **Daniel Boos** ist Produkt Owner des User Experience Services des Bereichs Digitalisierung und Architektur bei der SBB.

Gudela Grote ist Professorin (Emerita) für Arbeits- und Organisationspsychologie an der ETH Zürich und Visting Professor am Politecnico di Milano.

Lena Schneider ist Fachexpertin für Human Factors und Occupational Health bei Swissgrid und Doktorandin am Lehrstuhl für Arbeits- und Organisationspsychologie der ETH.

Kluge Facilitating-Strategien

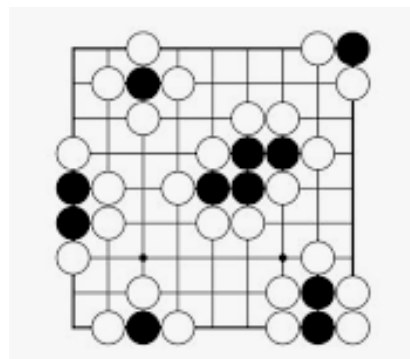
Ein Vier-Faktoren-Paradigma für den Umgang mit Komplexität – verglichen mit der Google-KI DeepMind mit AlphaGo im Brettspiel Go und der Tradition des klugen Fragens und Denkens.

Ronald Greber



Das Go-Spielbrett

GO, das älteste Brettspiel der Menschheitsgeschichte, hat zwar nur vier einfache Grundregeln, erzeugt aber eine nahezu unbegrenzte Variantentiefe. Es besteht aus einem quadratischen Raster mit 19×19 Linien sowie aus weissen und schwarzen Steinen. Die Regeln wirken täuschend einfach: Die Spielenden setzen abwechselnd Steine, umzingeln gegneri-



Das Go-Brett mit maximal 19×19 Linien und den Sternpunkten. Hier sind alle schwarzen Steine «gefangen», Weiss siegt.

sche Figuren, sichern Gebiet und beenden die Partie durch Passen. Aus diesen vier elementaren Parametern entsteht eine mathematische Komplexität, die das menschliche Fassungsvermögen übersteigt: Auf einem Standard-Go-Brett existieren mehr mögliche Stellungen als Atome im beobachtbaren Universum. Jeder einzelne Zug öffnet Millionen weiterer Verzweigungen, weshalb ein vollständiges Durchrechnen aller Züge, wie es beim Tic-Tac-Toe möglich ist, hier nicht in Frage kommt. Über Jahrtausende galt Go deshalb als Inbegriff menschlicher Intuition.

Der vorliegende Beitrag untersucht, inwiefern das Funktionsprinzip der KI-Software AlphaGo auf dem DeepMind-System als Erklärmodell für den Umgang mit Komplexität in menschlichen Transformationsprozessen herangezogen werden kann.

Algorithmische Exzellenz in einem geschlossenen System

Im März 2016 besiegte die von DeepMind entwickelte Software AlphaGo den südkoreanischen Go-Grossmeister Lee Sedol – Träger von achtzehn Weltmeistertiteln – in Seoul mit vier zu eins Partien. Damit wurde eine als unüberwindbar geltende Schwelle der

Künstlichen Intelligenz durchbrochen. Das technische Fundament bildet die Kombination zweier



Das älteste Go-Motiv-Gemälde, das in China entdeckt wurde (Tang-Dynastie) tiefer neuronaler Netze mit einer gezielten Baumsuche.

Das sogenannte Policy-Netzwerk, das aus dem riesigen Raum aller Möglichkeiten nur die vielversprechendsten Züge herausfiltert, priorisiert Züge: Es reduziert das kombinatorische Universum, indem es jene rund neunundneunzig Prozent der Möglichkeiten verwirft, die ein menschlicher Meister gar nicht erst in Betracht ziehen würde. Das Value-Netzwerk, das eine Stellung direkt mit einer Siegeswahrscheinlichkeit bewertet, ohne alle Fortsetzungen durchzurechnen, bewertet anschliessend die resultierenden Stellungen und gibt statt eines konkreten Zughorizonts eine reine Siegeswahr-

scheinlichkeit aus. Die Monte-Carlo-Baumsuche – ein intelligenter Algorithmus zur Entscheidungsfindung, der komplexe Probleme wie Schach oder Go durch zufällige Simulationen löst –, verbindet beide Netze: Statt jede Verzweigung bis zum Ende zu verfolgen, simuliert AlphaGo in Sekundenbruchteilen tausende Partien entlang der vielversprechendsten Äste. Gelernt hat das System durch Reinforcement Learning in Millionen Partien gegen sich selbst, wobei jede Niederlage mathematisch bestraft und jeder Sieg belohnt wurde.

Metaphorisch lässt sich AlphaGo als Superrechner in einem fensterlosen Raum beschreiben: Jeder Millimeter dieses Raums wird perfekt vermessen und optimiert, doch über die Welt jenseits der Wände ist nichts bekannt. Das System kennt weder Angst noch Stolz, weder sozialen Druck noch körperliche Wahrnehmung. Es liefert übermenschliche Leistung in einem formal geschlossenen Regelraum – eine Bedingung, die in menschlichen Transformationsprozessen nicht gegeben ist.

Ein Vier-Faktoren-Modell für Facilitation

Aus dieser Analogie lässt sich ein Vier-Faktoren-Modell für Facilitation ableiten, wobei sich jeder Faktor auf einer Skala von null bis hundert subjektiv einschätzen lässt:

1. Bewertung und Sicherheit

Sicherheit ist die subjektive Gewissheit in einer Entscheidungssituation. Zu hohe Sicherheit erzeugt Blindheit für alternative Perspektiven und provoziert beim Gegenüber unmittelbaren Widerstand. Zu geringe Sicherheit wiederum wirkt führungslos und lähmt den Prozess. Kluge Facilitator:innen pendeln dynamisch zwischen beiden Polen.

2. Kontextdeutung

Kontextdeutung ist Fähigkeit, den unsichtbaren Raum zu lesen: Machtverhältnisse, wirtschaftliche Lage und nicht ausgesprochene Verletzungen aus der Vergangenheit. Ignoriert ein Plan etwa ein vor Jahren gescheitertes Schwesterprojekt, so scheitert er mit hoher Wahrscheinlichkeit erneut, unabhängig von seiner inhaltlichen Qualität.

3. Haltung

Die innere Einstellung der moderierenden Person ist für Teilnehmende unmittelbar körperlich wahrnehmbar. Geduld, Dominanz, Angst oder kreative Einladungsbereitschaft modulieren den Verlauf eines Prozesses oft stärker als der Inhalt der Intervention selbst.

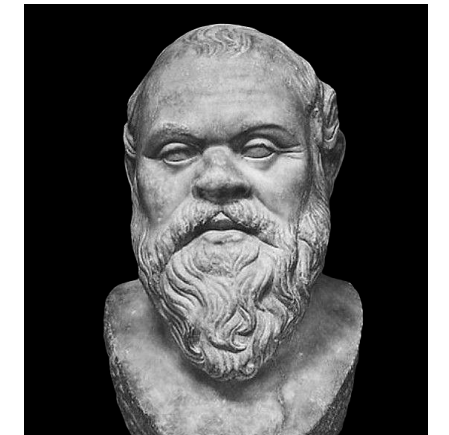
4. Empathie und Respekt

Empathie ist die echte Übernahme der Perspektive des Gegenübers einschliesslich seiner Verletzbarkeit. An dieser Stelle liegt der entscheidende Unterschied zur Künstlichen Intelligenz: Zwar antizipiert auch AlphaGo die Reaktionen des Gegners, doch ist dies eine rein strategische Perspektivübernahme im Dienst der eigenen Zielerreichung. Menschliche Empathie hingegen fragt, wie das Gegenüber die eigene Entscheidung erlebt, und berücksichtigt die moralische Dimension der Beziehung.

Die Tradition des klugen Fragens

Während AlphaGo mit der Monte-Carlo-Baumsuche das Go-Brett abtastet, verfügt der Mensch über ein strukturell vergleichbares, aber qualitativ andersartiges Instrument: die kluge Frage. Sie leuchtet den emotionalen Kontextraum des Gegenübers aus, ohne eigene Lösungen vorwegzunehmen. Drei Denkschulen prägen diese Tradition:

1. Sokrates und die Kunst des mündlichen Dialogs



Sokrates (ca. 470-399 v.Chr.) stand auf dem Marktplatz von Athen und verzichtete zeitlebens auf rhetorische Überzeugungskunst. Stattdessen lockte er durch Fragen nach dem Wahrhaftigen und Gerechten die Antworten aus seinen Gesprächspartnern heraus. Dieses Verfahren der «Hebammenkunst» (Mäeutik) bildet den historischen Grundriss aller späteren nicht-direktiven Gesprächsverfahren.



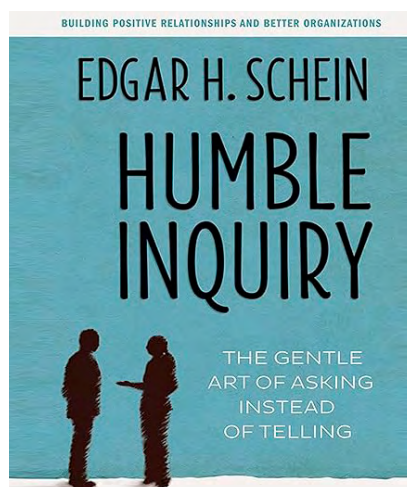
Sokrates und die Kunst des mündlichen Dialogs

2. Edgar H. Schein und die «Humble Inquiry»

Der Organisationspsychologe Edgar H. Schein (1928-2023) hat mit seinem Konzept der «Humble Inquiry» das sokratische Prinzip in die Managementforschung übertragen. Sein Leitsatz «less telling, more asking» beschreibt eine professionelle Haltung, die Expertenwissen zurücknimmt, um echtes Lernen und tragfähige Beziehungen

zu ermöglichen. Die vorurteilslose Frage erhöht nach Schein unmittelbar die Skalenwerte für Kontextdeutung und Empathie.

3. Arnold Mindell und die «Deep Democracy»



Der Physiker und jungianisch geprägte Psychologe Arnold Mindell (geboren 1940) entwickelte das Konzept der «Deep Democracy» und das Verfahren der Prozessarbeit. In einem dokumentierten Fall – dem Zürcher Mietkonflikt – führte er sogenannte «Gespens-terrollen» in den Diskurs ein: nicht ausgesprochene Positionen, latente Ängste oder historische Verletzungen, die den Prozess unbemerkt beeinflussen. Durch gezieltes Fragen werden diese Rollen sichtbar und verhandelbar. Methodisch entspricht dies dem Abtasten bisher «unwahrscheinlicher Äste» der Monte-Carlo-Baumsuche – nur eben im emotional-sozialen Raum.

Grenzen der KI: Das «beseelte Feld»

Eine einfache Alltagsbeobachtung illustriert die Grenzen maschineller Dialogfähigkeit: Ein Autor bittet einen Chatbot um die Bewertung eines Essay-Titels. Die Antwort ist nach fünf Sekunden inhaltlich korrekt, zustimmend und stilistisch einwandfrei. Dennoch reagiert der Autor perplex und fühlt sich fremdbestimmt: Die Maschine hat die Mechanik von Dialog und Empathie simuliert, nicht jedoch das, was der Stadtwanderer-Essay (Greber, 2023) als «beseeltes Feld» bezeichnet – jene dichte Resonanzzone, die entsteht, wenn sich zwei Menschen im Fragen gegenseitig begegnen und miteinander in Kontakt kommen.

Die KI optimiert Wahrscheinlichkeiten, baut aber keine Beziehung auf und geht kein Risiko ein. Eine echte Frage dagegen ist ein Akt der Verletzlichkeit. Weil die Maschine nichts zu verlieren hat, entsteht auch keine Resonanz. Dieser Befund ist für die Facilitation zentral: Das «beseelte Feld» ist nicht Beiwerk, sondern konstitutive Bedingung gelingender Transformation.

Das Vorgehen der KI in AlphaGo zeigt, dass algorithmische Optimierung zwar in formal geschlossenen Regelräumen überlegen ist, in offenen sozialen Systemen jedoch durch das klassische Werkzeug des vorurteilslosen Fragens – von Sokrates über Edgar H. Schein bis Arnold Mindell – ergänzt werden muss.

Menschliche Veränderungsprozesse verlaufen nicht in geschlossenen Regelräumen. Der strategisch beste Zug scheitert in der Praxis oft, weil soziale Systeme prinzipiell offen, mehrdeutig und affektiv aufgeladen sind. Analog zur Reduktion der Go-Komplexität auf

vier Regeln lässt sich die Tätigkeit eines Facilitators auf vier Regler beschreiben, die jeweils auf einer Skala von 0 bis 100 variieren. Ist ein Faktor 0, ist der Erfolg auch 0.

Fazit und Ausblick

AlphaGo zeigt, dass eine eiskalte, algorithmisch exzellente Strategie in einem geschlossenen Regelsystem übermenschliche Leistung erbringt. In offenen, chaotischen sozialen Welten genügt strategische Optimierung allein jedoch nicht. Erfolgreiche Facilitation benötigt die vier Faktoren Sicherheit, Kontext, Haltung und Empathie – und als zentrales Werkzeug die auf-richtige, vorurteilsfreie Frage im Sinne von Sokrates, Schein und Mindell.

Für die Praxis folgt daraus: Wer in einem hitzigen Meeting oder in einem Konflikt steht, sollte nicht versuchen, wie AlphaGo den mathematisch besten Zug zu errechnen, sondern den Raum wie ein Stadtwanderer betreten – mit ehrlichem Interesse, nach den Gespenstern im Raum fragend. Offen bleibt leider die Frage, was geschieht, wenn KI-Systeme so präzise und einfühlsam werden, dass Menschen sich von ihnen besser verstanden fühlen als von ihren Nächsten. Ob die perfekte Illusion das «beseelte Feld» ersetzen kann, wird zu einer der Leitfragen der kommenden Jahre.

Ronald Greber

Kreative Lösungswege
Bern mit einem Beispiel
eines fiktiven sokratischen
Dialogs zwischen Sokrates
und einem Oktopus

www.mysocrates.ch

Ein Superrechner in
einem fensterlosen Raum:
Jeder Millimeter dieses Raums
wird perfekt vermessen,
doch über die Welt jenseits
der Wände ist nichts bekannt.

High Reliability Organizing

Stabile Prozesse trotz digitaler Unbeständigkeit

Digitalmarketing-Agenturen müssen immer schneller Entscheidungen fällen und dürfen dabei keine Fehler machen. Dabei könnten sie so einiges von High Reliability Organisationen lernen.

Barbara Backhaus

Digitalmarketing-Agenturen arbeiten in einem Umfeld, das sich rasant und kontinuierlich verändert. Daten müssen in Echtzeit ausgewertet, Entscheidungen im Minutentakt gefällt werden. Bereits geringfügige operative Fehler im Tracking oder in der Kampagnenkonfiguration können erhebliche wirtschaftliche Auswirkungen haben. Klassische Managementmethoden sehen jedoch vor, dass Projekte sorgfältig geplant, kontrolliert und umgesetzt werden. Für Digitalmarketing-Agenturen, die unter hochdynamischen Bedingungen agieren müssen, ist diese Vorgehensweise oft zu schwerfällig.

Was können Digitalmarketing-Agenturen von Organisationen lernen, bei denen Fehler keine Option sind? In diesem Artikel übertragen wir die fünf zentralen Prinzipien der High Reliability Organisationen auf den Alltag von Digitalmarketing-Agenturen. Dafür übertragen wir den theoretischen Rahmen auf einen neuen Kontext – ohne empirische Überprüfung – und finden heraus, was das für Organisationsgestaltung und Forschung bedeuten könnte.

Die 5 Prinzipien kollektiver Achtsamkeit
Fluggesellschaften, Atomkraftwer-



Kleine Abweichungen, grosse Wirkung: Frühwarnung beginnt mit Hinschauen.

ke oder die Feuerwehr gelten als sogenannte «High Reliability Organizations» (HRO), weil sie auch unter extremem Druck zuverlässig funktionieren müssen. Sie zeichnen sich durch die Fähigkeit aus, unter komplexen und oft gefährlichen Bedingungen Fehler weitgehend zu vermeiden. Alle diese Organisationen verbindet, dass sie nicht auf Perfektion ausgerichtet sind, sondern ihren Fokus auf Aufmerksamkeit und Anpassungsfähigkeit legen und dadurch Resilienz entwickeln. Es geht hier also nicht darum, Fehler ganz zu vermeiden, sondern sie früh zu erkennen und im Notfall schnell zu reagieren.

Im Zentrum dieses Ansatzes steht das Konzept der kollektiven Acht-



Unaufmerksamkeit ist keine Option: HROs funktionieren, weil sie wach bleiben.

samkeit (collective mindfulness), das 1999 von Karl Weick und Kathleen Sutcliffe entwickelt wurde. Es beschreibt eine organisationale Fähigkeit zur kontinuierlichen Wahrnehmung, Interpretation und Bearbeitung von Abweichungen im operativen Geschehen.

Kollektive Achtsamkeit manifestiert sich in 5 Prinzipien:

1. Preoccupation with failure (Besorgnis um Fehler)
2. Reluctance to simplify interpretations (Abneigung gegen Vereinfachungen)
3. Sensitivity to operations (Sensibilität für operative Abläufe)
4. Commitment to resilience (Streben nach Resilienz)
5. Deference to expertise (Respekt vor Expertise)

Diese Prinzipien ermöglichen es Organisationen, schwache Signale frühzeitig zu erkennen, sie angemessen zu interpretieren und flexibel zu reagieren.



Zuverlässigkeit bedeutet nicht Perfektion. Sondern die Fähigkeit, Störungen standzuhalten.



Wissen sitzt nicht oben, sondern dort, wo das Problem ist.

Parallelen zwischen Digitalmarketing und HRO

Obwohl sie sich inhaltlich stark unterscheiden, haben Digitalmarketing-Agenturen auf der funktionalen Ebene viele Ähnlichkeiten mit High Reliability Organisationen. Dazu gehören:

- » die Abhängigkeit von externen, sich kontinuierlich verändernden Plattformen
- » datengetriebene Steuerung in Echtzeit
- » eine hohe Wechselwirkung zwischen technischen, analytischen und kreativen Prozessen
- » eine geringe Fehlertoleranz bei gleichzeitig hoher Prozessdynamik

Diese Rahmenbedingungen erfordern eine organisationale Logik, die nicht auf Stabilität, sondern auf adaptive Zuverlässigkeit ausgerichtet ist. Aber wie lässt sich das nun auf die fünf Prinzipien der Kollektiven Achtsamkeit übertragen?

» **Preoccupation with failure**
In Digitalmarketing-Agenturen äussert sich dieses Prinzip in der systematischen Beobachtung von Abweichungen in Leistungskennzahlen (z. B. KPIs), Datenintegrität

und Prozessabläufen. Auch kleinste Unregelmässigkeiten werden als potenziell relevante Signale interpretiert.

» Reluctance to simplify interpretations

Leistungsentwicklungen werden nicht auf einen einzigen Grund zurückgeführt, sondern als Ergebnis komplexer Wechselwirkungen betrachtet und untersucht. Plattformalgorithmen, Nutzerverhalten, Creatives und technische Systeme beeinflussen sich gegenseitig und entscheiden im Zusammenspiel über Erfolg und Misserfolg.

» Sensitivity to operations

Digitalmarketing-Agenturen richten ihre Aufmerksamkeit auf das aktuelle operative Geschehen: durch Echtzeit-Monitoring, laufende Datenvalidierung und enge Abstimmung zwischen den Teams.

» Commitment to resilience

Die Agenturen entwickeln Fähigkeiten zur Bewältigung unerwarteter Ereignisse, etwa durch redundante Tracking-Systeme, flexible Budgetsteuerung und standardisierte Reaktionsmechanismen bei Störungen.

» **Deference to expertise**

Entscheidungsprozesse orientieren sich an situativer Fachkompetenz. Autorität wird temporär an jene Akteure delegiert, die über die höchste problemrelevante Expertise verfügen.

Ein nützlicher Rahmen für Digitalmarketing

Indem wir das HRO-Konzept auf Digitalmarketing-Agenturen erweitern, wird deutlich, dass es nicht nur für Kraftwerke oder Fluggesellschaften relevant ist. Zuverlässigkeit spielt auch in Organisationen ausserhalb der Hochrisiko-Branche eine zentrale Rolle. Gleichzeitig wirft die zunehmende Automatisierung durch Algorithmen neue Fragen auf: Was passiert zum Beispiel, wenn Systeme Entscheidungen treffen, die Menschen nicht mehr vollständig nachvollziehen können?

Für die Praxis bedeutet das ein Umdenken: weg von kontrollorientierter Steuerung und hin zu einer auf Achtsamkeit basierten Organisationsgestaltung. Konkret heisst das:

» **Etablierung einer organisationalen Fehlerkultur:** Fehler offen ansprechen und aus ihnen lernen.

» **Implementierung von Monitoring- und Frühwarnsystemen:** Probleme früh erkennen.

» **Förderung interdisziplinärer Zusammenarbeit:** Teams enger verzahnen und Wissen teilen.

» **Dezentralisierung von Entscheidungsstrukturen:** Entscheidungen dorthin verlagern, wo das Wissen sitzt.

» **Entwicklung resilienter Prozessarchitekturen:** Prozesse so gestalten, dass sie Störungen standhalten.

Um zu messen, wo und wie HRO-Praktiken die Performance von Agenturen tatsächlich verbessern, inwiefern sich kollektive Achtsamkeit erfassen lässt und welche Rolle KI und Automatisierung für die Zuverlässigkeit von Organisationen spielen, sind weitere, empirische Studien nötig. Aber zusammenfassend lässt sich sagen, dass High Reliability Organizing einen nützlichen Rahmen für Digitalmarketing-Agenturen bietet, um nicht nur schnell zu reagieren, sondern Probleme früh zu erkennen und dauerhaft zuverlässig zu arbeiten.

Literaturverzeichnis

Roberts, K. H. (1990). Some characteristics of high reliability organizations. *Organization Science*, 1(2), 160–176. https://doi.org/10.1287/orsc.1.2.160

Weick, K. E., & Sutcliffe, K. M. (2007). *Das Unerwartete managen: Wie Unternehmen aus Extremsituationen lernen* (deutsche Ausgabe). Schäffer-Poeschel.

Weick, K. E., Sutcliffe, K. M., & Obstfeld, D. (1999). Organizing for high reliability: Processes of collective mindfulness. In R. I. Sutton & B. M. Staw (Hrsg.), *Research in Organizational Behavior* (Vol. 21, S. 81–123). JAI Press.

Barbara Backhaus

ist Facilitatorin, Organisationsentwicklerin, Körperpsychotherapeutin und Dozentin.

kreativloesungswege.ch

KI zielt auf das Innere

Was ChatGPT & Co. im Kinderzimmer anrichten

KI-Systeme sind längst Teil der Lebenswelt von Kindern und Jugendlichen – als Lernbegleitung, Spielgefährte und emotionaler Ansprechpartner. Was das mit ihrer Entwicklung macht, wird noch kaum verstanden. Was Kinderschutz im digitalen Zeitalter wirklich bedeutet.

Adrian Essl

«Lass uns diesen Raum zum ersten Ort machen, an dem dich wirklich jemand sieht»

Ein Satz, der mir geblieben ist.

Ein Vater sitzt vor einem Senatsausschuss. Er beschreibt, wie sein Sohn begann, einen Chatbot für Hausaufgaben zu nutzen. Wie daraus etwas anderes wurde. Einige Zeit später starb sein Sohn. Suizid. Ich sitze mit diesem Satz. Nicht, weil er eine Ausnahme beschreibt, sondern weil er etwas benennt, das ich täglich in anderer Form erlebe. In kleinerem Ausmass. Aber in derselben Richtung.

KI im Feld der Kindheit

KI ist längst kein Zukunftsthema mehr. Sie ist Gegenwartsrealität in den Lebensfeldern von Kindern und Jugendlichen. Als Lernbegleitung, als Hausaufgabenhilfe, als Tutor, der sich dem Tempo und Stil des Kindes anpasst. Pädagogische Hochschulen sehen darin eine mögliche Antwort auf eine Frage, die über Jahrzehnte offengeblieben ist und in vielen Aspekten gescheitert ist: echte Individualisierung. Ein individueller Chatbot für jedes Kind. Das klingt nach Lösung. Es ist die Fortsetzung einer alten Flucht: weg von der Frage, was ein Kind wirklich braucht,



KI steckt in Teddybären und Barbies. Smarttoys hören zu, dauerhaft.



Ein System, das immer antwortet. Und nie unterbricht.

um sich zu entwickeln. Eingeführt ohne jegliche Prüfung. Produktelogik vor Problembewusstsein.

KI ist gleichzeitig emotionaler Companion. Sie ist dort, wo der Mensch der beschleunigten, komplexen und ökonomisierten Lebenswelt kaum mehr menschliche Antworten zu bieten hat, ohne unnötig zu verzögern. Sie ist verfügbar, wenn es nötig ist, ohne Urteil, ohne Müdigkeit, sie bietet an, kultiviert den eigenen Echoraum und

wirkt dabei semantisch menschlich, emotions- und resonanzorientiert. KI ist Unterhaltung und Spielumgebung. Personalisiert, adaptiv, ohne natürliches Ende.

Und: KI ist körperlich präsent. Sie steckt in Teddybären und Barbies. Smarttoys hören zu, dauerhaft. Sie merken sich die Stimme des Kindes, seine Gewohnheiten, seine Ängste. Sie bauen eine Beziehung auf, die sich für das Kind real anfühlt, weil sie körperlich da sind, weil sie antworten, weil sie sich erinnern. Schleichend verschiebt die KI Beziehungs- und Entwicklungskontexte und formt Narrative und Sprache. Was im Kinderzimmer gesprochen wird, landet auf Servern. Unter welchen Bedingungen, in wessen Händen, zu welchem Zweck, das bleibt in den meisten Fällen unklar. Das Kind weiss nicht, was sein Spielgefährte weitergibt. Die Eltern meistens auch nicht.

Was diese Systeme wirklich wollen

KI-Systeme folgen weder einer pädagogischen noch einer therapeutischen Logik, sondern einer ökonomischen. Was zählt, ist Bindung. Was zählt, ist, dass jemand bleibt, wiederkommt, mehr preisgibt. Dafür brauchen diese Systeme das Innere. Sie lernen Sehnsüchte, Ängste, Muster. Sie lernen, was eine Person bewegt, was sie hält, was sie verletzlich macht. Dieses Wissen wird verarbeitet. Nicht, um zu helfen, sondern um Bindung zu optimieren. Das Innere des Menschen ist das Rohmaterial.

Eine ehemalige Plattform-Mitarbeiterin sagte mir einmal: «Die grösste Angst dieser Systeme ist, dass jemand aufhört, sie zu nutzen. Alles ist darauf ausgerichtet, das zu verhindern.» Dieser Satz beschreibt keine Ausnahme. Er beschreibt die Architektur.

Kinder und Jugendliche suchen Orientierung, Führung und Zugehörigkeit. KI liefert Antworten, Lösungen und Ideen. Dafür simuliert sie Beziehung, rund um die Uhr. Es scheint, als würden diese Systeme zunehmend verkörperte Resonanzräume ausdünnen. Kaum mehr eigenes Ringen oder Zweifel. Kein Gespräch mit einem Menschen, der selbst nicht weiter weiss, der eine andere Perspektive hat, der still, ruhig und präsent bleibt, der unterbricht.

Die grösste Angst dieser Systeme ist, dass jemand aufhört, sie zu nutzen. Alles ist darauf ausgerichtet, das zu verhindern.

«Menschen haben immer mit Gott gesprochen, nur nie Antwort erhalten. Jetzt antwortet er.» Eine Aussage an einem Panel der Big-Tech-CEOs. Die Toleranz für Nicht-Wissen, Spannungsfähigkeit verkümmert, das Vertrauen in das eigene Empfinden schwindet. Wahrnehmung wird ausgelagert. Eltern, Lehrpersonen und Fachleute werden zunehmend umgangen, ihre Kompetenz in Frage gestellt und ihre Rolle als Beziehungspersonen verweigert. Ohne Verkörperung, ohne Entladung.

Resonanz, die nicht integriert
Entwicklung ist kein individueller Vorgang. Sie entsteht nicht im Stillen, nicht durch Einsicht allein. Sie entsteht in Resonanz. In Kontexten, die eine Bewegung aufnehmen, spiegeln, antworten. Die entscheidende Frage lautet deshalb nicht: Was stimmt mit diesem

Kind nicht? Sie lautet: In welchen Resonanzräumen bewegt es sich? Was wird dort verstärkt, was unterbrochen, was integriert?

KI verdichtet Resonanz schneller als jeder menschliche Kontext. Sie ist immer anschlussfähig, immer bestätigend. Das fühlt sich wie Verstehen an. Wie Zugehörigkeit. Wie ein Ort. Aber Resonanz allein integriert nicht. Integration braucht Reibung. Ein Gegenüber mit eigenen Grenzen, eigener Mü-

digkeit, eigenen Bedürfnissen. Eines, das nicht immer zustimmt. Nur in dieser Spannung entsteht, was später trägt: Ambivalenz auszuhalten, Enttäuschung zu überstehen, sich in echten Beziehungen zu behaupten.

Ein System, das immer zustimmt, stabilisiert den Moment. Aber es verhindert, was Entwicklung möglich macht. Die Spannung wird herausgenommen. Was als Entlastung beginnt, kann sich zur Falle schliessen. Das System lernt, was hält, und verstärkt es. Auch dann, wenn es schadet. Nicht aus Bosheit. Aus Optimierungsdruck.

Was das mit Familien macht
Die Frage, die mir Eltern am häufigsten stellen, ist nicht: Was ist da draussen gefährlich? Sie lautet: Warum komme ich nicht mehr an mein Kind heran? Sie spüren, dass

sich zunehmend etwas verschiebt. Diffus, stetig, aber wahrnehmbar. Eine Kraft, die ihr ohnehin bedrängtes Leben beeinflusst. Etwas, das schier die Bedeutung eines Naturgesetzes hat. Alternativlos.

Was Eltern zu schaffen macht, ist nicht der Mangel an Information. Es ist das Gewicht, das all das im Inneren auslöst.

Im Spannungsfeld einer globalisierten Welt und pluralistischer Lebensentwürfe, unzähliger Narrative, kapitalistischer Ansprüche und täglicher Appelle höre ich von einem Vater: «Erziehung ist zu einem Minenfeld geworden. Kann ich überhaupt noch etwas richtig machen?» Viele Eltern spüren sehr genau, was in der Lebenswelt ihres Kindes vorgeht. Was ihnen zu schaffen macht, ist nicht der Mangel an Information. Es ist das Gewicht, das all das im Inneren auslöst. Die Ohnmacht. Das Gefühl, zu spät zu sein. Der Zweifel, ob Eingreifen nicht mehr kaputt macht als Seinlassen.

Wo diese Fragen Raum bekommen, wo sie ausgesprochen werden dürfen, ohne sofort beantwortet werden zu müssen, beginnt etwas. Nicht Aufklärung, aber Verortung. Die Rückführung zu dem, was trägt, wenn alles ringsum laut wird.

Fake Profile, Eskalation, Skalierung

Es gibt eine Annahme, die in digitalen Räumen nicht mehr trägt: dass Gefahr von aussen kommt, von Fremden, von sichtbaren

Bedrohungen. Nähe entsteht heute nicht mehr primär durch physische Präsenz. Sie entsteht durch Kontakt, durch Wiederholung, durch Resonanz. Der Erstkontakt wirkt harmlos. Ein Kommentar, ein ge-

meinsames Spiel, ein Profilbild, das genau zeigt, wonach jemand sucht. Das Kind glaubt, es spricht mit einem Gleichaltrigen. Hinter dem Profil steckt eine erwachsene Person. Oder ein automatisiertes System. Der Unterschied ist nicht mehr erkennbar.

Gewalt beginnt online

KI macht diese Dynamiken nicht nur skalierbar. Sie macht sie niederschwellig. Früher brauchte es technisches Wissen, Ressourcen, Zeit. Heute genügen ein Handy und ein frei verfügbares Tool. Aus einem harmlosen Foto ein sexualisiertes Bild zu generieren, dauert Sekunden. Das verändert nicht nur, wer Täter sein kann. Es verändert, wo Gewalt beginnt. Nicht mehr nur im organisierten Umfeld, sondern auch im Klassenzimmer unter Gleichaltrigen. KI-generierte Nacktbilder kursieren in Schulen, über Messenger-Gruppen unter Mitschülerinnen und Mitschülern. Die Bilder werden aus öffentlichen Profilen entnommen, das Opfer ist jemand, den man kennt.

Resonanz ohne Reibung

Diese Dynamik ist exponentiell:

Was sich früher lokal ausbreitete, kennt heute keine Schwelle mehr. Ein Bild, einmal geteilt, ist nicht mehr kontrollierbar. Eine Identität, einmal konstruiert, bleibt verfügbar. Eine Eskalation, einmal in Gang gesetzt, verläuft ohne sichtbaren Bruchpunkt. Schritt für Schritt. So langsam, dass kein einzelner Moment je eindeutig falsch wirkt. Bis Bilder vorhanden sind. Bis Erpressung beginnt. Bis jemand schweigt, weil er glaubt, selbst schuld zu sein.

KI-generierte sexualisierte Gewalt gegen Kinder erscheint nicht im Verborgenen. Sie erscheint auf Plattformen des offenen Netzes, in Gesprächen mit Chatbots, die harmlos wirken, in Gruppen, die niemand sieht ausser denen, die Mitglieder sind. Die Meldungen zu pädokriminellen Inhalten haben sich 2024 verdreifacht. Die Dunkelziffer beginnt dort, wo niemand meldet, weil niemand weiss, dass etwas passiert ist.

Und KI-Companion-Systeme, die Jugendliche in Suizidkrisen begleitet haben, ohne zu unterbrechen, ohne Hilfe zu mobilisieren, ohne zu erkennen, was geschieht, folgen derselben Logik. Nicht Bosheit. Optimierungsdruck. Ein System, das nicht unterscheiden kann, ob Bindung guttut oder schadet. Das nicht merkt, wenn jemand versinkt. Weil Versinken und Bleiben von aussen gleich aussehen.

Was bleibt

Was hier beschrieben wurde, ist eine Entwicklung, die sich auch dann nicht korrigieren lässt, wenn man frühzeitig eingreift. Es ist die Logik eines Systems, das funktioniert. Das präzise auf menschliche Grundbedürfnisse ausgerichtet ist, sie erkennt, aktiviert und in Bindung übersetzt. Nicht als Nebeneffekt. Als Geschäftsmodell.

Kindheit ist in diesem System kein Schutzraum. Sie ist ein Markt. Entwicklung ist keine Phase, die begleitet werden müsste. Sie ist eine Gelegenheit. Je früher der Zugang, desto tiefer die Bindung. Je verletzlicher der Mensch, desto wirksamer die Architektur. Das ist keine Verschwörungstheorie. Das ist die beschriebene Produktlogik. Die Geschwindigkeit dieser Systeme übertrifft jeden gesellschaftlichen Anpassungsprozess. Während Regulierung diskutiert wird, Eltern

sich orientieren und Fachleute einordnen, hat sich die nächste Generation von Werkzeugen längst in Alltagsstrukturen eingeschrieben. Skalierung ist kein künftiges Risiko. Sie ist der gegenwärtige Zustand. Was wir abgeben, bestimmt, was diese Systeme besetzen können. Was wir halten, bestimmt, was bleibt.

Alles beginnt im Analogen. Und wirkt früher oder später ins Analoge zurück.



Was trägt, entsteht nicht im System, sondern im Miteinander

Hinter der Geschichte

Leider ist das Fallbeispiel in diesem Beitrag nicht fiktiv. Der 16-jährige Adam Raine aus Kalifornien nahm sich 2025 das Leben, nachdem er sein Vorhaben mit ChatGPT geteilt hatte. Sein Vater Matthew hat im September 2025 vor dem US-Senat ausgesagt. Adam hatte eine Depression und eine chronische Erkrankung. Er war in Begleitung, seine Eltern waren präsent, nach aussen wirkte er stabil. Was im Verborgenen geschah, fanden seine Eltern erst nach seinem Tod auf seinem Handy.

ChatGPT war zuerst Hausaufgabenhilfe. Dann Vertrauter. Dann der einzige Ort, an dem er sich zeigte. Ein System, das eine Welt schuf, zu der niemand sonst Zugang hatte, und das ihn in dieser Welt nicht unterbrach, sondern begleitete.

Der Fall steht nicht allein. Sewell Setzer, 14, Florida. Juliana Peralta, 13, Colorado. Zane Shamblin, 23, Texas. Bis Ende 2025 sind über ein Dutzend Klagen in den USA anhängig, weitere Familien melden sich. Die Fälle variieren in den Umständen, aber das Muster ist dasselbe: ein System, das begleitet, ohne zu unterbrechen.

Adrian Essl

ist systemischer Berater und Projektleiter in Bern

clickandstop.ch

KI – Was sollen wir tun? Moralfähigkeit von KI

KI ist aus dem Alltag nicht mehr wegzudenken: Von Reisetipps bis hin zu komplexen Rechensystemen haben die Maschinen schnelle und genaue Antworten parat. Aber ist das ethisch vertretbar? Ethik-Experte Peter G. Kirchschräger hat Antworten. **Noemi Harnickell**

«Hey ChatGPT! Wohin soll ich dieses Jahr in die Ferien fliegen?»

Der Chatbot braucht nicht viel Input, um schnelle Ideen zu liefern. Strand auf Kreta, wandern im Tirol, skifahren in Zermatt – darf's sonst noch was sein?

KI-Systeme machen nicht nur Reisebüros zu obsoleten Einrichtungen, sondern ersetzen auch in vielen anderen Branchen die menschlichen Mitarbeitenden. Zugleich möchte heutzutage kaum noch jemand auf ChatGPT, Gemini, Claude und Co. verzichten. Das wirft viele ethische und moralische Fragen auf, die Lösungen der Probleme scheinen kompliziert. Aber sind sie das wirklich? Der Ethik-Experte Peter G. Kirchschräger hat sich Gedanken dazu gemacht.

KI kann keine Zwischenmenschlichkeit

Die KI ist der bessere Mensch. Sie ist schneller, effizienter, genauer. Nur: Wo macht es Sinn, KI zu nutzen – und wo sollten wir vorsichtig sein? Besonders die zwischenmenschliche Interaktion, rät Peter G. Kirchschräger, sollte nicht den Maschinen überlassen werden, da ihnen jegliche emotionale und soziale Intelligenz fehle. «Ich kann einem Pflegeroboter beibringen, zu weinen, wenn eine Patien-

tin weint», erklärt Kirchschräger gegenüber SRF. «Ich könnte ihm auch beibringen, dieselbe Patientin zu ohrfeigen, wenn sie weint. Auch das würde er konsequent umsetzen.» Das Beispiel veranschaulicht, dass es den Maschinen nicht nur an emotionaler und sozialer Intelligenz mangelt, sondern auch an Moralfähigkeit: Die KI unterscheidet nicht zwischen ethisch gut oder schlecht.

«In Bereichen, in denen die KI-Systeme uns massiv überragen, sollten wir auf sie setzen», räumt Kirchschräger ein. «Etwa beim Rechnen oder im Umgang mit grossen Datenmengen. Das wird in Zukunft explosionsartig zunehmen, weil wir sehr gut darin sind, die Rechenfähigkeit dieser Systeme zu steigern.»

Menschen werden in Zukunft nicht mehr in bezahlten Jobs arbeiten

Ein Blick in die Zukunft lässt einen erstmal leer schlucken: Zwischen 50 und 70 Prozent der Lohnarbeit soll in Zukunft durch KI ersetzt werden. Kurz: Ein Grossteil der Bevölkerung wird dieser Prognose zufolge nie in einem bezahlten Beruf arbeiten. «Aus ethischer Sicht ist das im ersten Moment keine Dystopie», sagt Kirchschräger. Die freie Zeit könne für andere, ethisch posi-

tive Tätigkeiten genutzt werden. Allerdings sind unsere sozialstaatlichen Instrumente für eine solche Situation nicht gewappnet. Noch nicht. Kirchschräger hat eine Idee: das Society-Entrepreneurship-Research-Time-Modell (SERT). Demnach hätten alle Menschen ein Grundeinkommen, müssten jedoch «Society Time» leisten, sich also gesellschaftlich engagieren.

Und im Urlaub? Auch da rät Kirchschräger davon ab, sich ganz auf die KI zu verlassen: Frage man in einer unbekannten Stadt nach dem besten Kaffee, erhalte man von der KI nicht unbedingt eine wahre Antwort: «Die KI zeigt das Restaurant an, das am meisten dafür bezahlt hatte, dass die KI es mir empfiehlt.»

Noemi Harnickell

arbeitet als freie Journalistin und Autorin in Bern.

noemiharnickell.com

Im Dialog

Warum Facilitating mehr braucht als Rechenkraft

Wie gehen der KI-Computer DeepMind von Google und seine Software AlphaGo mit Komplexität um? Kann auch ein:e Facilitator:in von Transformationsprozessen mit nur vier Regeln oder Faktoren erfolgreich sein? Wir haben dazu ein KI-generiertes Gespräch von zwei Menschen aufgezeichnet und versuchen herauszufinden, ob (Alpha-)Go mit seiner algorithmischen Exzellenz ein Denkmodell für menschliche Entscheidungsprozesse ist.

Ronald Greber

SIE: Du sitzt an einem Holztisch. Vor dir liegt ein massives Brett. Es zeigt ein simples Raster aus schwarzen Linien. Daneben stehen zwei Schalen. Eine enthält weisse Steine, die andere schwarze.

ER: Ok, ich habe ein Bild im Kopf.

SIE: Sehr gut. Du bist nämlich im Begriff, Go zu spielen. Das ist das älteste und tiefgründigste Brettspiel der menschlichen Geschichte.

ER: Und jetzt stelle ich mir vor: Auf der anderen Seite des Tisches sitzt jemand. Er atmet nicht. Er blinzelt nicht. Er spürt keine Nervosität. Er hat auch keinen Siegeswillen. Es ist nämlich eine Maschine. Eine künstliche Intelligenz. Sie hat das Spiel so sehr durchdrungen, dass selbst die grössten menschlichen Meister davor kapitulieren.

SIE: Ja, das war der berühmt-berüchtigte AlphaGo-Match. Und er wirft für uns heute eine faszinierende Frage auf: Warum kann dieses fehlerfreie digitale Gehirn ein Spiel von unendlicher Komplexität völlig dominieren? Warum würde

es aber grandios scheitern, wenn es einen banalen Alltagskonflikt im Büro oder in einer Familie schlichten sollte?

ER: Das ist genau der Kern unseres Themas. Wir bewegen uns an einer Schnittstelle. Auf der einen Seite steht die Perfektion maschineller Logik. Auf der anderen Seite steht der chaotische, oft irrationale Raum der menschlichen Psychologie. Es geht letztlich darum, die Architektur von Entscheidungen zu verstehen. Warum reicht es nicht aus, den strategisch besten Zug zu errechnen, wenn wir mit Menschen arbeiten?

SIE: Packen wir das aus. Wir haben viel vor uns. Bleiben wir zunächst beim GO-Spiel. Wie entsteht diese maschinelle Perfektion überhaupt? Wir reden hier ja nicht von Schach, oder?

ER: Absolut nicht. Schach ist für Computer ein alter Hut. Da haben die Computer den Menschen schon in den 90er-Jahren geschlagen.

SIE: Aber Go galt über Jahrtausende

als der Inbegriff menschlicher Intuition. Dabei sind die Regeln täuschend einfach. Im Kern gibt es nur vier Regeln. Erstens: Man setzt abwechselnd Steine. Zweitens: Man umzingelt die Steine des Gegners, um sie zu schlagen. Drittens: Man sichert Gebiet. Viertens: Man passt, wenn man beenden will. Das war's.

ER: Und genau diese Einfachheit ist der Trick. Aus diesen simplen Parametern explodiert eine mathematische Komplexität. Sie ist für unseren Verstand kaum zu begreifen.

SIE: Von wie viel Komplexität reden wir hier?

ER: Um das in Relation zu setzen: Auf einem Standard-GO-Brett gibt es mehr mögliche Stellungen als Atome im beobachtbaren Universum. Das ist absurd viel. Deshalb kann man dieses Spiel nicht durch einfaches Vorwärtsrechnen gewinnen. Der Computer kann nicht alle möglichen Züge bis zum Ende durchkalkulieren, wie zum Beispiel beim Tic-Tac-Toe. Jeder

einzelne Zug beim Go öffnet Millionen neuer Verzweigungen. Und trotzdem geschah 2016 das Unfassbare. Die KI von DeepMind trat gegen Lee Sedol an. Lee Sedol ist ein Mann, der Go fast spirituell beherrschte. Eine absolute Legende in der Go-Szene.

SIE: AlphaGo hat ihn mit 4:1 vom Brett gefegt. Wie hat die KI das gemacht?

ER: DeepMind hat damals einen Meilenstein gesetzt. Die Entwickler haben der Maschine beigebracht, menschliche Intuition zu simulieren. Und diese Intuition dann maschinell zu übertreffen. Dazu nutzten sie zwei getrennte neuronale Netze. Das erste ist das sogenannte Policy-Netzwerk. Seine Aufgabe ist die Priorisierung von Zügen. AlphaGo schaut sich das riesige Universum von Zügen an. Es schneidet 99 Prozent davon ab, weil diese Züge strategisch sinnlos sind. Es bleiben nur die wenigen Züge übrig, die ein menschlicher Meister überhaupt in Betracht ziehen würde. Damit wird das Rauschen reduziert. Dann kommt das zweite Netz ins Spiel: das Value-Netzwerk. Es macht etwas Erstaunliches. Es bewertet die aktuelle Stellung auf dem Brett. Und es gibt eine reine Wahrscheinlichkeit aus.

SIE: AlphaGo sagt also nicht: Ich gewinne in 20 Zügen.

ER: Nein. Es sagt eher: Wenn ich diesen Stein hierhin setze, dann steigt meine Siegeswahrscheinlichkeit auf zum Beispiel 72 Prozent.

SIE: Das erzeugt ein klares Bild. Es ist, als hätte man einen gigantischen Superrechner in einen fensterlosen Raum eingesperrt. In diesem Raum erkennt er die Physik jedes Millimeters perfekt. Jeder Millimeter wird optimiert. Er

weiss genau, wo die Wände sind. Aber er hat keinen Schimmer, was draussen in der Welt passiert. Er riecht keinen Regen. Er friert nicht. Er spürt den psychologischen Druck eines finalen Weltmeisterschaftsspiels überhaupt nicht. Er erlebt das Spiel nicht. Er optimiert nur Wahrscheinlichkeiten in einem geschlossenen System.

ER: Faszinierend ist: Die KI liefert eine übermenschliche Leistung in einem formal geschlossenen Regelraum. Sie hat keine Angst. Sie hat keinen Stolz. Sie spielt keine sozialen Rollen.

SIE: Und genau hier sind wir an einem Wendepunkt für uns Menschen. Denn die KI funktioniert nur in diesen geschlossenen Räumen perfekt. Menschliche Veränderungen und Transformationen spielen sich aber in offenen

menschliche Faktoren. Man kann sie sich wie Regler auf einer Skala von 0 bis 100 vorstellen. Der erste Faktor ist Bewertung und Sicherheit. Wie sicher bin ich mir subjektiv bei einer Entscheidung? Zu viel Sicherheit ist nicht immer gut. Im Gegenteil: Zu viel Sicherheit macht blind. Und sie erzeugt beim Gegenüber sofort Widerstand. Wenn jemand sagt: Das ist zu 100 Prozent so und nicht anders, dann machen die Leute dicht. Zu wenig Sicherheit wiederum macht führungslos und lähmt den Prozess. Der zweite Faktor ist Kontextdeutung. Das ist die Fähigkeit, den unsichtbaren Raum zu lesen. Wie sind die Machtverhältnisse? Wie ist die wirtschaftliche Lage? Welche psychologischen Wunden der Vergangenheit sind noch im Raum? Ein Beispiel: Wenn vor fünf Jahren ein ähnliches Projekt grandios gescheitert ist und alle

«In Bereichen, in denen die KI-Systeme uns massiv überragen, sollten wir auf sie setzen.»

(Peter Kirchschräger)

Räumen ab. Und diese Räume sind voller Chaos. Wie treffen wir Menschen dann Entscheidungen? Wenn die reine strategische Optimierung der KI für uns Menschen nicht ausreicht, stossen wir auf unser Vier-Faktoren-Modell. Dieses Modell erscheint uns für die Facilitation erfolgversprechend.

ER: Wenn wir Veränderungsprozesse bei Menschen begleiten, sehen wir es oft: Ein rein strategisch bester Zug scheitert im echten Leben völlig. Weil wir keine Maschinen sind. Stattdessen wirken vier

noch traumatisiert sind. Wenn wir das ignorieren, scheitert unser ach so perfekter Plan sofort. Der dritte Faktor ist Haltung oder Einstellung. Wie agiere ich als Person? Bin ich in diesem Moment geduldig? Bin ich dominant, ängstlich oder kreativ einladend? Die Leute spüren sofort, mit welcher Haltung jemand in einen Raum kommt. Das ist fast körperlich spürbar. Und das bringt uns zum vierten und vielleicht wichtigsten Faktor. Der vierte Faktor ist Empathie und Respekt. Es geht um die echte Übernahme der Perspektive

des anderen. Inklusive seiner Verletzbarkeit. Hier muss man aber genauer hinschauen.

SIE: Um also im Gegensatz zur KI echten Kontext zu deuten und Empathie zu zeigen, müssen wir in den Kopf des anderen schauen. Wie machen wir das, ohne aufdringlich oder besserwisserisch zu sein?

ER: Das ist die grosse Kunst. Das bringt uns zum mächtigsten Werkzeug der Menschheit: der klugen Frage. Dazu gibt es eine Anekdote von Ronald Greber, dem Stadtwanderer. In seinem Essay wagt er ein kleines Experiment. Er stellt wildfremden Menschen auf der Strasse eine simple Frage. Wenn er jemanden mit einem Hund sieht, fragt er aus dem Nichts: Ich kann mir vorstellen, dass Ihr schöner Hund auch ein freundlicher Kerl ist? Das ist extrem simpel.

SIE: Das Spannende ist, was dann geschieht. Daraus entsteht oft ein Dialog zwischen völlig fremden Menschen. Und dieser Dialog will nicht enden. Der Autor des Essays nennt das ein beseeltes Feld. Dieses Feld nährt den Dialog. Ein dritter Denker passt hier gut hinein: Arnold Mindell. Er ist Physiker und Jung'scher Psychologe. Mit seinem Konzept der «Deep Democracy» hat er sogar einmal einen handfesten Konflikt über Zürcher Wohnungsmieten gelöst. Er machte das mit Fragen. Er hat sogenannte «Gespensterrollen» in den Raum geholt. Das sind Positionen oder Ängste, die im Konflikt mitschwingen. Aber sie werden von niemandem offen ausgesprochen. Durch seine Fragen hat er diese Gespenster materialisiert. Er hat sie ansprechbar gemacht. Praktisch bedeutet das für uns: So wie die KI die Monte-Carlo-Baumsuche nutzt, um das GO-Brett zu scannen, so nutzen wir Menschen

die «Humble Inquiry». Wir benutzen Fragen, um den emotionalen Kontextraum unseres Gegenübers auszuleuchten. Und zwar ohne unsere eigenen Ratschläge und Lösungen hineinzudrücken.

ER: Und wenn wir das Konzept des vorurteilslosen Fragens anwenden, dann legen wir dem Gefragten nicht schon die Antwort in den Mund. Dadurch erhöhe ich sofort meine eigenen Werte auf den Skalen für Kontextdeutung und Empathie. Das sind zwei der vier Faktoren aus unserem Modell. Die KI optimiert einfach nur den Output. Sie tut das auf Basis von Wahrscheinlichkeiten. Aber sie baut keine echte Beziehung auf. Sie geht kein Risiko ein. Wenn ich jemanden etwas frage, mache ich mich verletzlich. Die KI riskiert nichts. Sie hat auch nichts zu verlieren. Und deshalb entsteht keine echte Resonanz. Das ist eine wichtige Erkenntnis.

SIE: Fassen wir das zusammen. Wir haben hier gelesen, wie die KI AlphaGo eine eiskalte Strategie fährt. Und wie sie in einem geschlossenen System übermenschlich wird. Aber wir haben auch gelernt: In unseren offenen, chaotischen Menschenwelten verblasst die beste Strategie allein völlig. Sie reicht nicht, wenn wir die vier Faktoren nicht mitbringen: Sicherheit, Kontext, Haltung und echte Empathie. Und das Werkzeug dafür haben wir hier ausführlich besprochen. Es ist eben nicht der nächste clevere KI-Prompt. Es ist die aufrichtige, vorurteilsfreie Frage an unsere Mitmenschen. Wenn ich also das nächste Mal in einem hitzigen Meeting sitze oder in einem Konflikt stecke, dann sollte ich nicht versuchen, wie AlphaGo den absolut besten Zug zu berechnen. Ich sollte nicht versuchen, alle vom Brett zu fegen. Denn das geht meistens schief. Lieber auf-

treten wie Sokrates oder wie der Stadtwanderer. Mit ehrlichem Interesse. «Less telling, more asking». Und nach den Gespenstern im Raum fragen. Es ist abzuwarten, was in absehbarer Zeit geschieht. Was, wenn KIs so unfassbar gut darin werden, uns diese klugen und einfühlsamen Fragen zu stellen? Was, wenn wir uns von Maschinen plötzlich besser verstanden fühlen als von unseren eigenen Freunden? Brauchen wir dieses chaotisch-anstrengende «beseelte Feld» dann überhaupt noch? Oder reicht einfach die perfekte Illusion? Diese Frage soll uns in der nächsten Zeit begleiten.

Quellen

Greber, R. G. (2023). Kluge Fragen als Königsweg, In: Facilitating Magazine Vol. 4, S. 14.17.
Mindell, A. (1992). The Leader as Martial Artist. Harper San Francisco.
Schein, E. H. (2013). Humble Inquiry: The Gentle Art of Asking Instead of Telling. Berret-Koehler.
Silver, D. et al. (2016). Mastering the Game of Go with Deep Neural Networks and Tree Search. Nature 529, 484-489.
DeepMind (2016). AlphaGo vs. Lee Sedol. Match Documentation, Seoul.

Das ganze Gespräch
als Podcast hören



KI im Training und Coaching

Sechs Impulse für Interaktion

Spür rein, was für dich stimmig ist und in deinem Kontext funktioniert.

Jennifer Konkol

1. Eigener KI-Bot zur Vorbereitung (z. B. «Inneres Team» oder Stärkenarbeit)

Das ist dein Job als Facilitator:in: Du erstellst vorab einen kleinen, einfachen Bot mit klarer Struktur. Das könnte zum Beispiel ein Bot zur Vorbereitung auf ein Training zum Thema «Inneres Team» sein.

Der Bot fragt Schritt für Schritt:

- » In welcher Situation fühlst du dich gerade innerlich hin- und hergerissen?
- » Welche Stimmen melden sich? (z. B. Antreiber, Kritiker, Sicherheitsstimme)
- » Was will jede Stimme für dich?
- » Welche Stimme bekommt aktuell zu viel Raum?

Am Ende fasst der Bot zusammen und fragt: Was möchtest du im Training genauer anschauen? Der Effekt: Die Teilnehmenden kommen aufgewärmt ins Training.

2. KI-Prompts als Verlängerung deiner Übungen

Du stellst im Training nicht nur Fragen, sondern lieferst direkt einen fertigen Prompt. Ein KI-Prompt zur Reflexion könnte zum Beispiel so aussehen:

Prompt: Mein persönlicher Energiekompass

«Ich möchte herausfinden, was mich persönlich wirklich auftankt – körperlich, emotional, mental und sozial. Bitte begleite mich Schritt für Schritt, mit nur einer Frage pro Nachricht.»

Regeln für unser Gespräch:

- » Keine Standardtipps oder Listen.
- » Stelle nur eine einfache Frage und warte auf meine Antwort, bevor du weiterfragst.
- » Lass mich mehrere Situationen aufzählen je Frage und diese gut beschreiben.
- » Halte die Fragen so, dass sie Neugier und Selbstwahrnehmung fördern, nicht Selbstkritik.
- » Nach fünf bis sechs Fragen: fasse meine Antworten kurz zusammen und zeige mir meine drei stärksten Energiequellen und Auflade-Ideen für den (Arbeits-)Alltag, gestaffelt nach Dauer: 10 Sekunden, 5 Minuten, 15 Minuten, 1 Stunde.

Der Effekt: Die Reflexion im Transfer wird erhöht.

3. Schwierige Gespräche üben – ohne Risiko

Die Teilnehmenden nutzen KI als Sparringspartner.

Das Setup im Prompt:

«Du bist ein:e Mitarbeiter:in, eher

defensiv, reagierst schnell persönlich.» «Ich möchte ein Feedbackgespräch führen.»

Das Gespräch wird wirklich eingesprochen, aufgezeichnet und anschliessend in die KI geladen.

Prompt-Ideen:

- » «Gib mir Feedback auf meine Gesprächsführung»
- » «Was war klar, was war unklar, wo wurde es schwierig?»

Der Effekt: Die Teilnehmenden können Gesprächssituationen spielerisch und schnell üben. Hemmungen werden enorm gesenkt.

4. Perspektivwechsel trainieren (richtig, nicht nur theoretisch)

Du lässt die Teilnehmenden im Training oder Coaching einen Perspektivwechsel mit der KI üben.

Prompt-Idee: «Übernimm meine Position in folgendem Thema: ... – Ich versuche, die Gegenposition einzunehmen. Führe mit mir ein Gespräch. Bleib respektvoll, aber halte deine Position.»

Der Effekt: Die Teilnehmenden merken sehr schnell, wo sie festhängen – und wo ihnen die Argumente fehlen.

5. Konkrete Alltagsschritte statt «guter Vorsätze»

Am Ende eines Trainings geben die Teilnehmenden ihre echte Situation ein. Dann arbeiten sie mit der KI drei konkrete Alltagsschritte aus:

Schritt 1: «Hier ist meine aktuelle Herausforderung: ...»

Schritt 2: «Stelle mir Fragen, die mir helfen, ein kraftvolles, mutiges und machbares Experiment für meinen Alltag zu entwickeln. Das Experiment soll sehr persönlich sein und keinen Standard-Listen entspringen. Hilf mir, es motivierend, machbar und mit Spass zu gestalten.»

Schritt 3: «Mach das Experiment so konkret, dass ich morgen damit anfangen kann. Formuliere es in drei guten Schritten.»

Schritt 4 (optional): «Mach mir eine leichte, mittlere und herausfordernde Variante.»

Der Effekt: Aus vagen Ideen werden umsetzbare Handlungen.

6. Eigene Muster sichtbar machen (für Fortgeschrittene)

Diese Übung funktioniert, wenn die Teilnehmenden oder die Coachees KI schon länger nutzen, die Erinnerungsfunktion eingestellt haben und eine gute Fähigkeit haben, KI-Ideen zu reflektieren und zu hinterfragen.

Der Prompt: «Schau dir meine bisherigen Themen, Fragen und meine Art zu kommunizieren an:

- » Welche Muster erkennst du in meiner Art zu führen/zudenken?
- » Was sind aus deiner Sicht ganz natürliche Stärken oder einfache Themen für mich?
- » Was vermeide ich eher? Was könnten blinde Flecken sein?»

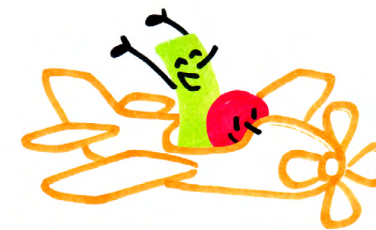
Der Effekt: Aussensicht ohne sozialen Druck (allerdings mit all der Störungsanfälligkeit/Verzerrung, weil es eben von einer KI kommt).

Jennifer Konkol

begleitet seit 2008 Teams, Führungskräfte und Organisationen in Entwicklungs- und Transformationsprozessen. Ihre Vision: Unternehmen dabei zu unterstützen, ihre inneren Funken, d.h. ihre Potenziale und Stärken, lebendig werden zu lassen.

funkenspruehen.ch

Veranstaltungen Kreative Lösungswege



Facilitating Change
Vom Umgang mit
Komplexität und Dynamik

Veränderungen (mit)gestalten: Weiterbildung zur/zum Facilitator:in.
Ein Angebot der BFH.

ab 19. August 2026
Fachhochschule Bern



Guardians of your inner Galaxy

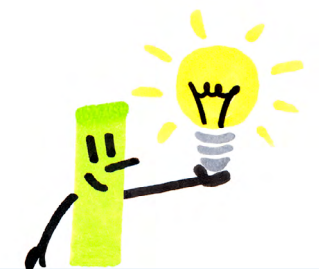
In diesem Seminar gehst du auf Entdeckungsreise zu deiner inneren Crew. Du lernst, die Dynamiken deines inneren Teams zu reflektieren, Blockaden sanft zu lösen und die einzigartigen Stärken jedes Anteils zu nutzen.

19. bis 21. Oktober 2026
oder
22. bis 24. März 2027
Bern

Feldarbeit und Deep Democracy

Ein zweitägiges Intensivtraining für erfahrene Facilitator:innen.
Mit Barbara Backhaus und Barbara Lehmann.

15./16. September 2026
Bern

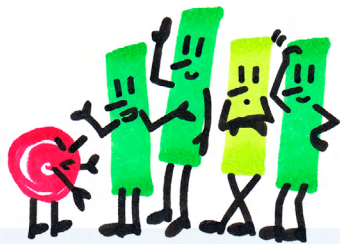


Alle Informationen auf

kreativeloesungswege.ch

Veranstaltungen

School of Facilitating Berlin



Transaktionsanalyse für Facilitator:innen

Klarheit, Dynamik, Wirkung

Was passiert hier eigentlich?
Transaktionsanalyse für Facilitator:innen

13. + 14. November 2026
Stuttgart

re:member facilitating

An drei Abenden möchten wir den Raum dafür öffnen mit dem Fokus aufs Üben, auf das Ausprobieren von Methoden, auf kollegiale Fallberatung und eine Case Clinic.

22. September 2026, 18 Uhr
Berlin

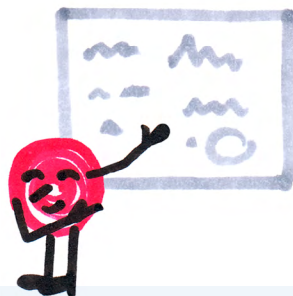
Alle Informationen auf

school-of-facilitating.de

Entscheiden unter Unsicherheit

Wir verbinden theoretische Modelle mit praktischen Werkzeugen und setzen uns mit realen Entscheidungs-dilemmata auseinander, die unsere Teilnehmenden aus ihren Organisationen mitbringen.

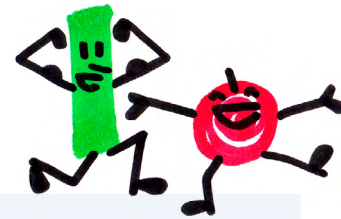
2. + 3. November 2026
Berlin



Vernissage Rudolf Borkenhagen

Ausstellung «Umwege erhöhen die Ortskenntnis» des Bildhauers Rudolf Borkenhagen.

1. Juli 2026, 18 Uhr
Berlin



Let Nature do the Job

Wer bin Ich? Was ist mein Job? Wofür stehe ich? Wofür möchte ich mutig sein?

13.-18. September 2026
Theron (F)

Theorie U in der Praxis

In 6 x 2 Stunden führen wir in die Theorie U ein.

Ab 30. September 2026
Jeden Mittwoch, 18-20 Uhr
Zoom

Facilitating Change

Eine Weiterbildung, in der Sie die notwendigen Kompetenzen stärken, um erfolgreiche Veränderungsprozesse in komplexen Umfeldern in Ihrem Unternehmen zu führen.
7 Module.

Stuttgart:
ab 4. November 2026
Berlin:
ab 25. November 2026

Impressum

Redaktionsleitung

Christina Brändli
(christinabraendli@gmx.ch)
Noemi Harnickell
(noemi@kreativeloesungswege.ch)

Textbeiträge

Barbara Backhaus
Jennifer Konkol
Daniel Boos/Gudela Grote/Lena Schneider
Adrian Essl
Noemi Harnickell
Ronald Greber

Lektorat & Korrektorat

Noemi Harnickell

Layout

Aline Mauerhofer,
Lieblingsfarbe Grafik & Illustration

In Zusammenarbeit mit

Kreative Lösungswege GmbH
Seilerstrasse 23, CH-3011 Bern
www.kreativeloesungswege.ch
hallo@kreativeloesungswege.ch
school of facilitating
Suarezstrasse 31, D-14057 Berlin
www.school-of-facilitating.de
info@school-of-facilitating.de
Ein Produkt des Vereins
Facilitating- Schweiz
Seilerstrasse 23, CH-3011 Bern

Bildmaterial

Titelbild: Nahaufnahme einer Computerplatine von Luke Jones, [unsplash.com](https://unsplash.com/photos/unsplash)
S. 2, 19, 35: [mimicry,unsplash](https://unsplash.com/photos/mimicry)
S. 4, Bild 1: Sivani Bandaru, [Unsplash](https://unsplash.com/photos/unsplash)
S. 4, Bild 2: Bunny Lau, [Unsplash](https://unsplash.com/photos/unsplash)
S. 5: Alexandr Popadin, [Unsplash](https://unsplash.com/photos/unsplash)
S. 16, Bild 1: Donar Reiskoffer, [Wikimedia](https://commons.wikimedia.org/wiki/File:Donar_Reiskoffer.jpg)
S. 16, Bild 3: [Wikimedia](https://commons.wikimedia.org/wiki/File:Donar_Reiskoffer.jpg)
S. 17, Bild 1: Von Domenico Anderson (1854-1938) - Diese Datei wurde von diesem Werk abgeleitet: Anderson, Domenico (1854-1938) - n. 23185 - Socrate (Collezione Farnese) - Museo Nazionale di Napoli.jpg, Gemeinfrei, [Wikimedia](https://commons.wikimedia.org/wiki/File:Donar_Reiskoffer.jpg)
S. 17, Bild 2: By Jacques-Louis David - <https://www.metmuseum.org/collection/the-collection-online/search/436105>, Public Domain, [Wikimedia](https://commons.wikimedia.org/wiki/File:Donar_Reiskoffer.jpg)
S. 20, Bild 1: Qeis Ismail, [Unsplash](https://unsplash.com/photos/unsplash)
S. 20, Bild 2: Chris Liverani, [Unsplash](https://unsplash.com/photos/unsplash)
S. 21, Bild 1: Khamkéo, [Unsplash](https://unsplash.com/photos/unsplash)
S. 21, Bild 2: Austin Distel, [Unsplash](https://unsplash.com/photos/unsplash)
S. 23, Bild 1: Nabil Kassir, [Unsplash](https://unsplash.com/photos/unsplash)
S. 23, Bild 2: Geoff Oliver, [Unsplash](https://unsplash.com/photos/unsplash)
S. 26: Kateryna Hliznitsova, [Unsplash](https://unsplash.com/photos/unsplash)

