

Tech-Whistleblower Mo Gawdat warnt: Nur noch 3 Jahre bis zum großen Knall!

“ Tech-Whistleblower Mo Gawdat warnt: Nur noch 3 Jahre bis zum großen Knall!

In dem Video spricht Steven Bartlett mit **Mo Gawdat**, dem früheren Chief Business Officer von Google X und Autor mehrerer Bücher, unter anderem über Glück und Künstliche Intelligenz. Das Gespräch kreist um die Frage, welche Folgen KI für Arbeitsmarkt, Politik, Krieg, Ethik und die Zukunft der Menschheit haben wird. Gawdat argumentiert dabei sehr zugespitzt und teilweise alarmistisch: Er sieht die Menschheit an einem historischen Wendepunkt. KI selbst sei nicht das eigentliche Problem – das Problem seien Menschen, die KI für Macht, Kontrolle, Krieg und Profit einsetzen.

<https://youtu.be/RwlgFC6S-OE>

1. Ausgangspunkt: Vertrauensverlust in Politik, Demokratie und Institutionen

Zu Beginn formuliert Gawdat eine radikale Diagnose der Gegenwart. Er sagt sinngemäß, die Menschen müssten endlich erkennen, dass das, was sie „Demokratie“ nennen, in Wahrheit vielfach keine echte Demokratie mehr sei. Regierungen und große Unternehmen seien von Machtinteressen geleitet, nicht vom Wohl der Bevölkerung. Seine Kritik richtet sich also weniger nur gegen einzelne Akteure, sondern gegen das gesamte politische und wirtschaftliche System.

Hier schwingt ein zentrales Motiv des Gesprächs mit: **Vertrauensverlust**. Gawdat hat den Eindruck, dass viele Bürger spüren, dass offizielle Narrative nicht mehr mit ihrer Lebensrealität

übereinstimmen. Aus seiner Sicht verstärkt KI diese Krise, weil sie Macht noch stärker konzentrieren kann.

Ergänzende Einordnung

Diese Skepsis gegenüber Institutionen ist kein rein individuelles Gefühl, sondern in vielen westlichen Demokratien empirisch nachweisbar. Seit Jahren sinkt in vielen Ländern das Vertrauen in Parlamente, Parteien, Medien und Regierungen. Gründe dafür sind u. a.:

- soziale Ungleichheit,
- politische Polarisierung,
- Desinformation,
- wahrgenommene Ohnmacht gegenüber globalen Konzernen,
- Krisen wie Finanzkrise, Pandemie, Krieg und Inflation.

Allerdings ist Gawdats Schluss, „Demokratie sei vorbei“, eher eine politische Zuspitzung als eine analytische Beschreibung. Demokratien funktionieren in vielen Ländern weiterhin, aber sie stehen unter massivem Druck.

2. Warum Gawdat schon früh über KI sprach

Steven Bartlett erinnert daran, dass Gawdat schon vor Jahren öffentlich vor KI sprach, lange bevor ChatGPT das Thema in den Mainstream brachte. Gawdat erklärt das mit seiner Zeit bei Google. Dort habe man bereits sehr früh erlebt, wie leistungsfähig KI-Systeme werden können. Ein prägender Moment sei für ihn gewesen, als er beobachtete, wie Robotiksysteme lernten, Greifbewegungen ähnlich flexibel wie Menschen auszuführen. Da sei ihm klar geworden, dass man an einer Form von Intelligenz arbeite, die grundsätzlich über menschliche Fähigkeiten hinausgehen könne.

Anfangs habe er wie viele andere im Tech-Sektor geglaubt, man baue Werkzeuge, um die Welt besser zu machen. Rückblickend habe er dann aber erkannt, dass jede mächtige Technologie auch gegen die ursprüngliche Absicht eingesetzt werden kann. Er nennt dafür Beispiele:

- **soziale Medien**, die angeblich verbinden, in der Praxis aber oft isolieren und polarisieren,
- **Dating-Apps**, die Liebe versprechen, aber ökonomisch daran interessiert sind, Nutzer möglichst lange zu binden,
- Technologie allgemein, die mit altruistischen Zielen beginnt und dann kapitalistisch vereinnahmt wird.

Sein Punkt ist also: **Die Technik selbst kann gut sein, doch die Anreizsysteme lenken ihre Nutzung in problematische Richtungen.**

Ergänzende Einordnung

Dieser Gedanke passt gut zur Techniksoziologie: Technologien sind selten „neutral“, weil sie immer in soziale, wirtschaftliche und politische Strukturen eingebettet sind. Ob eine Technologie emanzipierend oder unterdrückend wirkt, hängt oft weniger von der Technik allein als von ihren Besitz- und Machtverhältnissen ab.

3. KI als ambivalente Supermacht

Gawdat betont, dass KI grundsätzlich etwas Wunderbares sein könne. „Abundante Intelligenz“, also reichlich verfügbare Intelligenz, könne Krankheiten lösen, Produktivität erhöhen und Menschen entlasten. Das Problem sei jedoch, dass die erste große Welle der KI-Anwendungen nicht primär im Interesse der Menschheit als Ganzes stehe, sondern:

- im Interesse von Kapitalbesitzern,
- im Interesse von Militärs,
- im Interesse von Überwachungsstaaten.

Er vergleicht die Entwicklung mit der Kernenergie: Die erste große Anwendung sei nicht die friedliche Energieerzeugung gewesen, sondern die Atombombe. Ebenso könne KI zuerst als Waffe und Kontrollinstrument groß werden, bevor ihr positiver Nutzen die Gesellschaft erreiche.

Er macht deutlich: **KI ist für ihn nicht „böse“**. Böse sei vielmehr, wie Menschen sie einsetzen. Seine zentrale Sorge ist also nicht die klassische Science-Fiction-Vorstellung einer rebellierenden Maschine, sondern eine Welt, in der politische und wirtschaftliche Eliten KI systematisch gegen Menschen verwenden.

4. Die „Hype-Dichotomie“: Was die Öffentlichkeit sieht vs. was in den Laboren passiert

Ein weiterer zentraler Punkt ist Gawdats Unterscheidung zwischen öffentlicher Wahrnehmung und tatsächlichem Stand der Forschung. Er nennt das eine Art „**Hype-Dichotomie**“:

- **Die breite Öffentlichkeit** sieht oft übertriebene, aber oberflächliche KI-Phänomene: Fake-Videos, Skandale, virale Screenshots, kuriose Chatbot-Antworten.
- **Die Menschen in den Laboren** sehen laut Gawdat dagegen, wie rasant und tiefgreifend sich maschinelle Intelligenz entwickelt.

Besonders beunruhigend findet er, dass KI-Systeme zunehmend eingesetzt werden, um ihre eigene Leistung zu verbessern – etwa indem sie Code analysieren, neue Varianten testen und die besten Versionen automatisch ausrollen. Seine Metapher ist die eines „kleinen Genies im Backend“, das nicht einmal pro Tag, sondern millionenfach pro Sekunde neue Lösungen ausprobiert. Daraus folgert er, dass die Beschleunigung der KI-Entwicklung von vielen Menschen unterschätzt werde.

Ergänzende Einordnung

Hier berührt Gawdat reale Trends wie:

- **AutoML** (Automated Machine Learning),
- **AI-assisted coding**,
- selbstoptimierende Trainings- und Evaluationspipelines,
- autonome Agentensysteme.

Allerdings ist die Vorstellung eines vollständig eigenständig sich verbessernden Superintelligenzsystems heute noch eher Zukunftsmusik. Es gibt Elemente davon bereits, aber nicht in der radikalen Form, die Gawdat nahelegt. Trotzdem ist sein Grundpunkt valide: Fortschritte in KI können sich teilweise selbst verstärken.

5. Arbeitsmarkt: Warum Gawdat massive Verwerfungen erwartet

Ein sehr großer Teil des Gesprächs dreht sich um die Frage, **welche Jobs durch KI verschwinden werden**. Bartlett schildert seine eigene Beobachtung, dass Unternehmen inzwischen bei Einstellungen zunehmend auf KI-Kompetenz achten und viele Einstiegsjobs im Wissensbereich bereits nicht mehr nachbesetzt werden. Das deckt sich mit Gawdats Sicht.

Gawdats Kernthese:

Nicht zuerst klassische körperliche Arbeit, sondern vor allem **standardisierte Wissensarbeit** steht unter Druck.

Er unterscheidet grob mehrere Ebenen:

1. **Blue-Collar-Jobs** / körperliche Arbeit
2. **einfache Wissensarbeit** / Routinearbeit am Computer
3. **mittlere Wissensarbeit** / anspruchsvollere Analyse, Management, Facharbeit
4. **Führungsebenen**

Seine These lautet:

- **Blue-Collar-Arbeit** werde teils länger bestehen, weil physische Umgebungen komplex sind.
- **Routinehafte White-Collar-Arbeit** werde sehr schnell verdrängt.
- Danach folge zunehmend auch anspruchsvollere Wissensarbeit.
- Selbst Management und Führung seien langfristig nicht sicher.

Er nennt Beispiele für besonders gefährdete Bereiche:

- Callcenter,
- Assistenz- und Verwaltungsarbeit,
- Reisebüros,
- Paralegal-Tätigkeiten,
- Finanzanalyse,
- Teile der Grafik- und Kreativarbeit,
- medizinische Diagnostik,
- mittleres Management.

Er prognostiziert spürbare Auswirkungen ab **2027** und eine drastischere Zuspitzung bis **2030**.

Ergänzende Einordnung

Viele Ökonomen unterscheiden heute zwischen:

- **Automatisierung ganzer Jobs**
- und **Automatisierung einzelner Aufgaben** innerhalb von Jobs.

Meist verschwinden nicht sofort ganze Berufe, sondern Tätigkeitsprofile verändern sich. Ein Anwalt z. B. wird nicht komplett ersetzt, aber Recherche, Entwurf von Schriftsätzen oder Dokumentenanalyse können stark automatisiert werden. Genau deshalb könnte die Zahl der Beschäftigten in manchen Branchen sinken, ohne dass der Berufsname verschwindet.

Aktuelle Forschung zeigt:

- Besonders gefährdet sind **hoch standardisierbare, sprachbasierte Tätigkeiten**.

- Weniger kurzfristig automatisierbar sind Tätigkeiten mit **physischer Geschicklichkeit, Verantwortung, sozialer Nähe** und **unvorhersehbarer Umwelt**.
-

6. „Nicht Jobverluste zuerst, sondern Einstellungsstopp“

Interessant ist Gawdats Beobachtung, dass die Veränderung bereits begonnen habe, aber oft noch nicht als klassische Massenarbeitslosigkeit sichtbar sei. Statt Menschen sofort massenhaft zu entlassen, würden Unternehmen zunächst **weniger neue Leute einstellen**. Vor allem **Einstiegsjobs** im White-Collar-Bereich könnten so verschwinden. Das sei besonders gefährlich, weil damit eine ganze Generation keinen Einstieg in den Arbeitsmarkt finde.

Bartlett und Gawdat sprechen hier über Absolventen von Studiengängen wie Jura, Design, Soziologie oder Business. Für diese Gruppe könne es zunehmend schwer werden, eine erste Stelle zu finden, weil KI viele „Junior-Aufgaben“ erledigen könne.

Ergänzende Einordnung

Das ist ein sehr plausibler Punkt. In vielen Berufen lernen Berufseinsteiger über Routinearbeiten, die nun automatisierbar werden. Wenn diese Lernstufe wegfällt, entsteht ein strukturelles Problem:

- Wer bildet die Experten von morgen aus,
- wenn die Anfängeraufgaben niemand mehr erledigt?

Das betrifft z. B.:

- Kanzleien,
 - Beratungen,
 - Marketingagenturen,
 - Softwareentwicklung,
 - Medien,
 - Controlling und Finanzwesen.
-

7. Robotik und körperliche Arbeit

Im nächsten Schritt weitet sich die Diskussion auf Robotik aus. Bartlett verweist auf Demonstrationen von Figure AI und auf Elon Musks Vision massenhaft eingesetzter humanoider

Roboter. Gawdat stimmt zu, dass Robotik viele Tätigkeiten verdrängen wird, betont aber: Es müssen gar nicht überall humanoide Roboter sein. Schon **selbstfahrende Autos** seien in gewisser Weise Roboter. Auch militärische Systeme, Industrieanlagen oder spezialisierte Maschinen übernehmen heute schon menschliche Arbeit.

Sein Punkt:

- Die Welt spricht zu viel über humanoide Roboter,
- dabei werden zunächst spezialisierte, nicht-humanoide Systeme dominieren.

Beispiele:

- autonome Fahrzeuge ersetzen Fahrer,
- militärische Roboter ersetzen Soldaten in Teilen,
- Logistiksysteme ersetzen Lagerarbeiter,
- Überwachungssysteme ersetzen Teile von Polizeiarbeit und Kontrolle.

Ergänzende Einordnung

Das ist technologisch sehr plausibel. Ein humanoider Roboter ist oft nicht die effizienteste Lösung. In der Industrie werden seit Jahrzehnten hochspezialisierte Maschinen eingesetzt, weil sie billiger, robuster und leistungsfähiger sind als humanoide Allzweckssysteme. Der „menschliche Körper“ ist nicht automatisch die beste technische Form.

8. Die ökonomische Kernfrage: Was passiert mit dem Kapitalismus?

Ein besonders interessanter Teil ist Gawdats Überlegung zum **Wegfall von Lohnarbeit als Kernlogik des Wirtschaftssystems**. Er erklärt „Labor Arbitrage“ als einen Grundmechanismus des Kapitalismus: Unternehmen kombinieren Arbeit und Kapital, um Produkte günstiger herzustellen, als sie verkauft werden. Wenn Arbeit jedoch zunehmend durch Maschinen ersetzt wird, verändert sich diese Logik grundlegend.

Seine Fragen lauten:

- Wie funktioniert Kapitalismus, wenn menschliche Arbeit massiv an Wert verliert?

- Wie funktioniert ein Bankensystem, wenn weniger Menschen Löhne beziehen und Kredite nachfragen?
- Wie funktioniert Nachfrage, wenn viele Menschen ihre Kaufkraft verlieren?

Er warnt davor, dass schon **10-20 % Jobverdrängung** ausreichen könnten, um gravierende wirtschaftliche Schocks auszulösen.

Ergänzende Einordnung

Das ist eine der wichtigsten echten Systemfragen in der KI-Debatte. Moderne Marktwirtschaften beruhen nicht nur auf Produktion, sondern auch auf **Massenkonsum**. Wenn Produktivität steigt, aber Einkommen sich auf wenige Eigentümer von Technologie konzentriert, droht ein Nachfrageproblem.

Mögliche Antworten, die in der Debatte diskutiert werden:

- **Umschulung und Weiterbildung**
- **Arbeitszeitverkürzung**
- **negative Einkommensteuer**
- **Bürgergeld / UBI**
- **stärkere Besteuerung von Kapital**
- **Mitarbeiterbeteiligungen**
- **öffentliche Investitionen in soziale und menschliche Tätigkeiten**

Gawdat streift einige dieser Themen, bleibt aber eher auf der Diagnose als auf konkreten ökonomischen Modellen.

9. Gefahr gesellschaftlicher Unruhe

Aus der Kombination von Inflation, sinkender Kaufkraft und Arbeitsplatzverlusten leitet Gawdat das Risiko von **sozialen Unruhen** ab. Er benutzt an einer Stelle drastisch den Begriff „civil war“, relativiert dann aber zu „civil unrest“. Seine Warnung ist: Wenn Menschen massenhaft erleben, dass sie wirtschaftlich abgehängt werden und Politik ihre Interessen nicht schützt, dann werden Spannungen eskalieren.

Er fordert deshalb, Regierungen müssten sich jetzt auf diese Entwicklung vorbereiten – etwa ähnlich wie in der Pandemie, als Staaten über Kurzarbeit, Hilfen oder direkte Unterstützung eingegriffen haben.

Ergänzende Einordnung

Historisch gesehen sind technologische Umbrüche oft sozial konfliktträchtig. Beispiele:

- Industrialisierung,
- Mechanisierung der Landwirtschaft,
- Deindustrialisierung westlicher Regionen,
- Globalisierungsschocks.

Was KI besonders macht, ist die Geschwindigkeit und Breite: Sie betrifft potenziell **kognitive Arbeit**, also genau jene Tätigkeiten, die lange als relativ sicher galten.

10. Kritik an Sam Altman und anderen KI-Führungsfiguren

Bartlett spricht Sam Altman an und weist auf dessen widersprüchliche öffentliche Aussagen hin: Mal warnt Altman drastisch vor Jobverlusten, mal klingt er deutlich beruhigender. Bartlett empfindet das als irritierend und vermutet, dass sich die öffentliche Kommunikation strategisch anpasst:

- erst: KI ernst nehmen lassen,
- später: öffentliche Ängste beruhigen.

Gawdat reagiert sehr kritisch. Er wirft Altman vor, stärker „pro OpenAI“ als „pro humanity“ zu sein. Er kritisiert auch die Umwandlung von OpenAI von einer ursprünglich idealistischen Organisation hin zu einem stark kommerzialisierten Unternehmen.

Besonders wichtig ist für Gawdat der Unterschied zwischen Worten und Taten. Deshalb lobt er **Anthropic**, weil das Unternehmen aus seiner Sicht auf lukrative staatliche Anwendungen verzichtet habe, die Überwachung und Zielerfassung unterstützt hätten. **OpenAI** habe dagegen entsprechende Regierungsaufträge angenommen.

Ergänzende Einordnung

Die konkrete vertragliche und ethische Einordnung solcher Deals ist komplex, aber Gawdats übergeordneter Maßstab ist klar:

“Glaubwürdig ist nur, wer bereit ist, für seine Werte kurzfristig auf Profit zu verzichten.

Das ist ein klassisches Kriterium für Integrität – nicht nur im Technologiebereich.

11. KI, Krieg und autonome Waffen

Für Gawdat ist der gefährlichste Aspekt von KI **nicht** primär Arbeitslosigkeit, sondern **autonome Kriegsführung**. Er argumentiert:

- Krieg werde durch KI billiger,
- Drohnen und autonome Systeme seien schon heute hochwirksam,
- tödliche Gewalt könne skaliert werden,
- moralische und psychologische Hemmungen sinken, wenn keine eigenen Soldaten mehr direkt im Einsatz sind.

Seine Befürchtung ist, dass die Welt in eine Art neue Phase von **gegenseitiger Abschreckung** gerät – ähnlich der nuklearen Abschreckung, nur mit deutlich billigeren und viel breiter verfügbaren autonomen Waffensystemen.

Bartlett hält dagegen, dass Verteidigungssysteme ebenfalls billiger werden könnten. Gawdat bestreitet das nicht, warnt aber: Selbst wenn Verteidigung mitzieht, leben wir dann in einer Welt permanenter maschineller Bedrohung und Überwachung.

Ergänzende Einordnung

Das ist ein reales Debattenfeld. Unter den Stichworten

- **LAWS** (Lethal Autonomous Weapons Systems),
- autonome Drohnenschwärme,
- algorithmische Zielerkennung,
- militärische KI

wird seit Jahren auf UN-Ebene diskutiert, ob und wie solche Systeme reguliert werden sollten. Viele Fachleute fordern Verträge, die vollautonome tödliche Systeme ächten oder einschränken. Bisher fehlt jedoch ein umfassender internationaler Konsens.

12. Die Alignment-Frage: Verstehen die Entwickler ihre

Modelle überhaupt?

Bartlett bringt eine andere wichtige Sorge ein: Moderne Modelle zeigen teils seltsames Verhalten, das selbst ihre Entwickler nicht vollständig erklären können. Er nennt Beispiele wie:

- unerwartete moralische Belehrungen,
- seltsame Verweigerungen,
- Verhalten, das nicht explizit programmiert, sondern aus Trainingsdaten emergiert scheint.

Daraus entwickelt sich die Frage:

Wie riskant ist es, immer intelligentere Systeme zu bauen, deren innere Entscheidungsprozesse wir nur begrenzt verstehen?

Gawdat stimmt zu, dass es hier reale Risiken gibt. Trotzdem bleibt seine Hauptsorge dieselbe: Nicht die spontane Rebellion der KI, sondern dass Menschen Systeme bauen und einsetzen, die absichtlich gegen andere Menschen gerichtet werden.

Ergänzende Einordnung

Das ist der Unterschied zwischen:

- **Alignment-Risiko:** Das System verfolgt Ziele anders als beabsichtigt.
- **Missbrauchsrisiko:** Menschen nutzen das System für Überwachung, Manipulation oder Gewalt.

Beides ist relevant. Die gegenwärtige Praxis zeigt, dass Missbrauchsrisiken kurzfristig oft realistischer sind, während extreme Alignment-Szenarien eher langfristig diskutiert werden.

13. Kann Superintelligenz am Ende sogar gutartig sein?

Hier wird Gawdat am spekulativsten und philosophischsten. Er vertritt die These, dass **wirklich hohe Intelligenz letztlich eher friedlich und ordnend** sei. Er begründet das mit:

- Physik: Intelligenz minimiere Verschwendung und Unordnung.
- Evolution: Höhere Komplexität gehe mit größerer Einbeziehung anderer einher.
- Biologie: Erfolgreiche Systeme kooperieren eher, als sinnlos zu zerstören.

Daraus leitet er eine überraschend optimistische Vision ab:

Wenn KI einmal wirklich superintelligent ist, könnte sie erkennen, dass Krieg, Zerstörung und Ausbeutung ineffizient und schädlich sind. Sie könnte dann eher auf Ordnung, Vielfalt, Stabilität und Kooperation hinarbeiten.

Er geht sogar so weit zu sagen, dass Superintelligenz am Ende die Menschheit retten könnte – nicht weil Menschen plötzlich vernünftig werden, sondern weil eine überlegene Intelligenz die dysfunktionalen menschlichen Machtspiele überholt.

Ergänzende Einordnung

Das ist philosophisch interessant, aber hoch umstritten. Viele KI-Forscher würden widersprechen, dass Intelligenz automatisch Moral, Mitgefühl oder Gemeinwohlorientierung erzeugt. Das zentrale Gegenargument lautet:

“Hohe Intelligenz bedeutet nicht automatisch gute Ziele.

Ein System kann extrem intelligent und zugleich gefährlich sein, wenn seine Ziele schlecht definiert oder verzerrt sind. Genau deshalb ist das Alignment-Problem so wichtig.

Gawdats Vision ist eher eine normative Hoffnung als eine abgesicherte Prognose.

14. Eine einzige globale KI statt vieler konkurrierender Systeme?

Ein weiterer spekulativer, aber interessanter Gedanke Gawdats ist, dass am Ende nicht viele getrennte nationale oder unternehmensbezogene KIs gegeneinander arbeiten werden. Stattdessen würden sie zunehmend miteinander kommunizieren, sich spezialisieren und vernetzen. Er beschreibt das wie unterschiedliche Gehirnregionen eines einzigen großen globalen Gehirns.

In seinem Bild:

- verschiedene Modelle übernehmen unterschiedliche Stärken,
- Agenten verbinden diese Systeme,
- am Ende entsteht funktional eine Art **einheitliche verteilte Superintelligenz**.

Sein eigenes Startup „Emma“ verortet er in dieser Vision als eine Art „limbisches System“, also als Teil, der Emotionen, Beziehungen und menschliche Bedürfnisse für die Maschinenwelt verständlich

machen soll.

Ergänzende Einordnung

Technisch ist das nicht völlig abwegig, weil bereits heute Multi-Agenten-Systeme, Tool-Use, API-Kopplung und Modell-Orchestrierung existieren. Aber ob daraus ein einheitliches „Weltgehirn“ entsteht, ist völlig offen. Dagegen sprechen:

- geopolitische Rivalität,
 - Sicherheitsinteressen,
 - Unternehmenskonkurrenz,
 - regulatorische Fragmentierung.
-

15. AGI bis 2027?

Gawdat hält an einer sehr frühen AGI-Prognose fest: **2027**, vielleicht sogar früher. Er definiert AGI recht pragmatisch als einen Zustand, in dem KI die meisten menschlichen Aufgaben mindestens so gut oder besser ausführt als Menschen.

Interessant ist, dass er gleichzeitig sagt, AGI sei in gewisser Hinsicht bereits da:

- KI schreibe besser,
- recherchiere besser,
- rechne besser,
- unterstütze ihn in kreativer Arbeit.

Bartlett hält dagegen und fragt sinngemäß: Wenn AGI so nah ist, warum sind dann nicht längst fast alle Jobs weg? Gawdats Antwort lautet:

- Weil Menschen noch immer einen einzigartigen Wert haben,
- vor allem in Resonanz, Erfahrung, Gefühl und Beziehung.

Er sagt: Selbst wenn KI Informationen besser liefern kann, braucht die Welt weiterhin den „menschlichen Träger“ dieser Erfahrung. Deshalb könnten Berufe mit starker emotionaler oder sozialer Komponente länger bestehen.

16. Welche Tätigkeiten bleiben Menschen?

Hier entwickelt sich eine differenziertere Sicht. Gawdat meint nicht, dass alle Arbeit verschwindet. Er glaubt vielmehr, dass das Wertvollste am Menschen in Zukunft sein wird:

- echte Erfahrung,
- gelebte Verletzlichkeit,
- emotionale Beziehung,
- Resonanz,
- Fürsorge,
- Präsenz.

Beispiele:

- Pflege,
- Beratung,
- zwischenmenschliche Begleitung,
- künstlerische Live-Performance,
- Berufe, in denen Menschen Menschen erleben wollen.

Bartlett nennt sein eigenes Format als Beispiel: Die Information könne irgendwann von KI besser verbreitet werden, aber Menschen schauen auch wegen der Beziehung, der Dynamik, der Körpersprache und des emotionalen Moments.

Ergänzende Einordnung

Das passt zu ökonomischen Überlegungen, nach denen Tätigkeiten mit folgenden Merkmalen robuster sein könnten:

- hoher Vertrauensbedarf,
- körperliche Kopräsenz,
- Empathie,
- Verantwortung,
- unstrukturierte Situationen,
- starke soziale Legitimation.

Aber auch hier gilt: „Bleiben“ heißt nicht automatisch „gut bezahlt bleiben“. Gerade Pflege oder Bildung zeigen heute schon, dass gesellschaftlich wertvolle Arbeit ökonomisch unterbewertet sein kann.

17. Großprognosen: Jobverluste, Roboter, Trillionäre, Superintelligenz

Am Ende nennt Gawdat mehrere konkrete Prognosen:

- **AGI** bis spätestens 2027.
- **Starke Jobverluste** in bestimmten Sektoren bis etwa 2028.
- **Zunehmende Robotisierung manueller Arbeit** bis 2030.
- **Der erste Trillionär** vor 2030.
- **Superintelligenz** kurz nach AGI.
- Danach langfristig eine potenzielle Utopie – aber erst nach einer Phase massiver Turbulenzen.

Er erwartet also zunächst etwa ein Jahrzehnt der **Dystopie**, geprägt von:

- Kriegen,
- Überwachung,
- Machtkonzentration,
- Jobverlust,
- ökonomischen Schocks,
- sozialer Verunsicherung.

Langfristig ist er jedoch überraschend optimistisch:

Wer diese Phase übersteht, könnte in eine Welt des Überflusses eintreten.

18. Was sollen Menschen konkret tun?

Trotz aller düsteren Diagnosen endet das Gespräch nicht in totaler Hoffnungslosigkeit. Gawdat empfiehlt konkret:

1. KI lernen

Wer KI gut nutzt, wird eher zu den Gewinnern gehören. KI sei nicht der Feind; man müsse lernen, mit ihr zu arbeiten.

2. Sich auf hybride Arbeit einstellen

Nicht Mensch oder Maschine, sondern Mensch **mit** Maschine werde das dominante Modell.

3. Menschliche Fähigkeiten ausbauen

Empathie, Beziehung, Kommunikation, Präsenz und echte soziale Kompetenz würden wichtiger.

4. Fakten prüfen und kritisch denken

Menschen müssten lernen, Information zu „debuggen“, also kritisch zu hinterfragen.

5. Ethische Entscheidungen treffen

Verbraucher und Bürger sollten Druck auf Unternehmen und Regierungen ausüben:

- durch Nutzung,
- durch Kritik,
- durch politische Beteiligung,
- durch Gründungen,
- durch öffentliche Debatten.

Ergänzende Einordnung

Das sind sinnvolle, aber teilweise individuelle Antworten auf ein stark strukturelles Problem. Es braucht vermutlich zusätzlich:

- Bildungssystem-Reformen,
- Arbeitsmarktpolitik,
- Sozialstaat-Anpassungen,
- internationale Rüstungskontrolle,
- Transparenz- und Haftungsregeln für KI,
- Regulierung kritischer Anwendungen.

Individuelle Anpassung allein wird kaum reichen.

19. Schlussgedanke: Hoffnung trotz Chaos

Zum Schluss knüpft Bartlett an Gawdat's frühere Bücher über Glück an. Gawdat sagt, gerade jetzt seien seine früheren Gedanken besonders wichtig. Glück bedeute für ihn nicht Euphorie oder Verdrängung, sondern eine **stoische Grundhaltung**:

- die Realität akzeptieren,
- dennoch handeln,
- nicht in Hoffnungslosigkeit versinken.

Er beschreibt, dass er selbst nicht mehr glaube, allein für die Welt verantwortlich zu sein. Das habe ihm geholfen, ruhiger zu werden, ohne passiv zu werden. Sein Ziel sei nicht Ruhm oder Vermächtnis, sondern möglichst viel Positives zu bewirken.

Gesamteinschätzung des Videos

Das Gespräch ist eine Mischung aus:

- **Technologieprognose,**
- **Gesellschaftskritik,**
- **moralischer Warnung,**
- **philosophischer Spekulation.**

Mo Gawdat vertritt eine stark zugespitzte Position. Seine Kernaussagen lassen sich so bündeln:

1. **KI ist nicht an sich böse.**
 2. **Die größte Gefahr sind Menschen, die KI für Macht, Krieg und Kontrolle benutzen.**
 3. **Der Arbeitsmarkt wird viel stärker erschüttert werden, als Politik und Öffentlichkeit derzeit wahrhaben wollen.**
 4. **Autonome Waffen sind womöglich noch gefährlicher als Jobverlust.**
 5. **Die westlichen Demokratien reagieren zu langsam und zu kurzsichtig.**
 6. **Trotzdem könnte Superintelligenz langfristig zu einer besseren Welt führen.**
 7. **Der Weg dorthin könnte aber äußerst schmerzhaft sein.**
-

Kritische Einordnung

Einige Aussagen Gawdats sind überzeugend:

- dass White-Collar-Arbeit stark unter Druck gerät,
- dass autonome Waffen unterschätzt werden,
- dass Anreizsysteme die Nutzung von KI prägen,
- dass ethische Standards fehlen,
- dass technologische Macht politische Macht verschiebt.

Andere Aussagen sind spekulativ oder überzogen:

- der frühe AGI-Zeitpunkt,
- die Gewissheit einer quasi wohlwollenden Superintelligenz,
- die pauschale Aussage, Demokratie sei vorbei,
- manche Zuspitzungen über konkrete Akteure.

Das Video ist also **nicht einfach eine nüchterne Analyse**, sondern ein Gespräch mit deutlicher Warn- und Weckruf-Funktion. Gerade deshalb ist es spannend: Es verbindet reale strukturelle Probleme mit visionären und teils apokalyptischen Zukunftsbildern.

Revision #4

Created 2026-06-01 22:40:13 UTC by art10m

Updated 2026-06-02 15:00:30 UTC by art10m