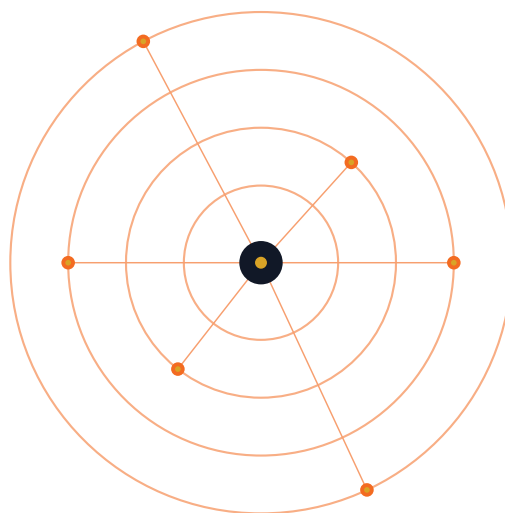


ADL

Architektur des Lebens

Mathematische Grundlegung eines mehrskaligen relationalen Ordnungsmodells

Gerichtete Wirkungen, Referenzrahmen, Rückkopplung, Driftzerlegung und menschlich kontrollierte Intervention



Elemente – Beziehungen – Wirkungen – Zeit – Rückkopplung

M. Selman Yildizhan

Epoch Creation, Hamburg, Germany

14. Juli 2026

Projekt-DOI: [10.5281/zenodo.21209954](https://doi.org/10.5281/zenodo.21209954)

Kurzfassung

ADL wird als diskretes, mehrskaliges und rückgekoppeltes System auf einem gewichteten gerichteten Graphen formalisiert. Das Modell trennt vier Ebenen, die in informellen Beschreibungen häufig vermischt werden: den beobachteten oder geschätzten Systemzustand, die modellierte Beziehungsstruktur, den gesetzten Bewertungsrahmen und die menschlich verantwortete Intervention. Eine dimensionslose Abweichungsfunktion verbindet Zustands-, Wirkungs-, Grenz- und Instabilitätsterme. Ein daraus abgeleiteter Balancewert ist ausschließlich ein relativer Modellindex.

Die Fassung 2.0 ergänzt eine exakte Driftzerlegung. Damit lässt sich unterscheiden, ob sich das System verändert hat, ob lediglich die angenommene Beziehungsstruktur aktualisiert wurde oder ob der Bewertungsrahmen verschoben wurde. Für die Verdichtung stabiler Teilordnungen wird eine prüfbare Coarse-Graining-Abbildung mit Aggregationsfehler eingeführt. Ein vollständig gerechnetes Beispiel, Minimalannahmen, Widerlegungskriterien und eine Reproduzierbarkeitsmatrix begrenzen den Geltungsanspruch.

Status der Arbeit

Dieses Dokument legt eine intern konsistente Modellarchitektur und ein reproduzierbares Rechenbeispiel vor. Es ist weder ein empirischer Nachweis universeller Gültigkeit noch ein Beweis einer neuen Naturgesetzlichkeit. Ob ADL in einem konkreten Feld beschreibt, prognostiziert oder Interventionen verbessert, muss mit vorab definierten Daten, Baselines, Alternativmodellen und wiederholbaren Tests geprüft werden.

Abstract

ADL is formalized as a discrete, multiscale, feedback-driven system on a weighted directed graph. The framework separates four layers that are often conflated in informal system descriptions: the observed or estimated system state, the modeled relational structure, the declared evaluation frame, and the human-authorized intervention. A dimensionless deviation functional combines state, effect, constraint, and temporal-instability terms; its monotone transformation is a relative model score, not an absolute measure of truth, justice, or equilibrium.

Version 2.0 introduces an exact drift decomposition that distinguishes system movement from relation-model updates and changes of the evaluation frame. Stable subgraphs may be coarse-grained only when an explicit aggregation map and its information-loss residual are reported. The paper provides a worked numerical example, minimal well-posedness conditions, falsification criteria, and a reproducibility protocol. It establishes a formal architecture, not empirical universality.

Schlüsselbegriffe

ADL; Architektur des Lebens; gewichteter gerichteter Graph; dynamisches System; Referenzrahmen; Rückkopplung; Driftzerlegung; projizierte Dynamik; Coarse-Graining; Human-in-the-loop.

Zitationsvorschlag

Yildizhan, M. Selman (2026): *ADL – Mathematische Grundlegung eines mehrskaligen relationalen Ordnungsmodells*. Working Paper, Version 2.0. Epoch Creation, Hamburg. Projekt-DOI: [10.5281/zenodo.21209954](https://doi.org/10.5281/zenodo.21209954).

Beitrag und Positionierung

Was diese Fassung formal leistet

1. Sie definiert Zustände, Beziehungen, Wirkungen, Referenzen, Grenzen und Eingriffe als getrennte mathematische Objekte.
2. Sie verhindert durch explizite Skalierung, dass Größen mit unvereinbaren Einheiten unbemerkt addiert werden.
3. Sie zerlegt beobachtete Drift in Systembewegung, Beziehungs- beziehungsweise Modellupdate und Referenzrahmenverschiebung.
4. Sie trennt rechnerische Vorschläge von menschlicher Auswahl und protokolliert Unsicherheit sowie Interventionskosten.
5. Sie macht Mehrskaligkeit durch eine Aggregations- und Rekonstruktionsabbildung samt Verlustmaß prüfbar.
6. Sie ordnet den Erkenntnisstand über eine Anspruchsleiter von interner Konsistenz bis zur Übertragbarkeit ein.

Gewichtete Graphen, Verlustfunktionen, projizierte Dynamik, Systemidentifikation, Sensitivitätsanalyse, Coarse-Graining und Human-in-the-loop-Entscheidungen sind etablierte mathematische beziehungsweise methodische Bausteine. Der Beitrag dieser Arbeit liegt in ihrer transparenten Zusammenführung zu einer einheitlichen ADL-Modellarchitektur, in der expliziten Rollen- und Ebenentrennung sowie in der hier angegebenen Driftbuchhaltung. Eine Prioritäts- oder Neuheitsbehauptung gegenüber sämtlicher Fachliteratur wird nicht erhoben.

Begriffskontrolle

Zentrum bedeutet Referenzvektor, nicht automatisch Mittelwert, Gleichgewicht oder Wahrheit. **Balance** bezeichnet den Wert einer offengelegten Bewertungsfunktion, nicht Gerechtigkeit an sich. **Korrektur** bedeutet eine Intervention, die unter einem festgehaltenen Rahmen einen kleineren prognostizierten Wert erzeugt; daraus folgt keine moralische Richtigkeit. **Beziehung** bezeichnet zunächst eine modellierte gerichtete Kopplung; Kausalität entsteht nicht allein durch einen Eintrag in der Matrix.

Inhaltsverzeichnis

1	Problemstellung und Erkenntnisvertrag	3	10 Prüfung, Falsifizierbarkeit und Anspruchsleiter	10
2	Mathematische Grundobjekte	3	11 Identifizierbarkeit, Unsicherheit und Grenzen	11
3	Wirkung auf dem gerichteten Graphen	4	12 Rohform für den Kopf	12
4	Dimensionslose Abweichungsfunktion	5	A Symboltabelle	13
5	Dynamik, Rückkopplung und Lernen	6	B ADL-Kernsystem	14
6	Drift ohne verschobene Torpfosten	7	C Reproduzierbarkeitsprotokoll	15
7	Menschlich kontrollierte Intervention	7	D Literatur und methodische Anschlüsse	16
8	Rekursion und Mehrskaligkeit	8	E Autor und Projekt	16
9	Vollständig gerechnetes Beispiel	9		

1. Problemstellung und Erkenntnisvertrag

Viele Modelle erfassen entweder Zustände, Netzwerke, zeitliche Dynamik oder Optimierungsziele. ADL verbindet diese Ebenen in einem Prüfkreislauf: Was wurde beobachtet? Welche Beziehung wird angenommen? Welche Wirkung folgt daraus? Gegen welchen offengelegten Rahmen wird sie bewertet? Welche Änderung kommt mit der Zeit zurück? Wer entscheidet über einen realen Eingriff?

Kernfrage

Wie verändert sich ein relationales System unter modellierten Eigen- und Fremdwirkungen, und welcher Anteil einer beobachteten Abweichungsänderung stammt aus dem System, aus dem Beziehungsmodell oder aus einem veränderten Bewertungsrahmen?

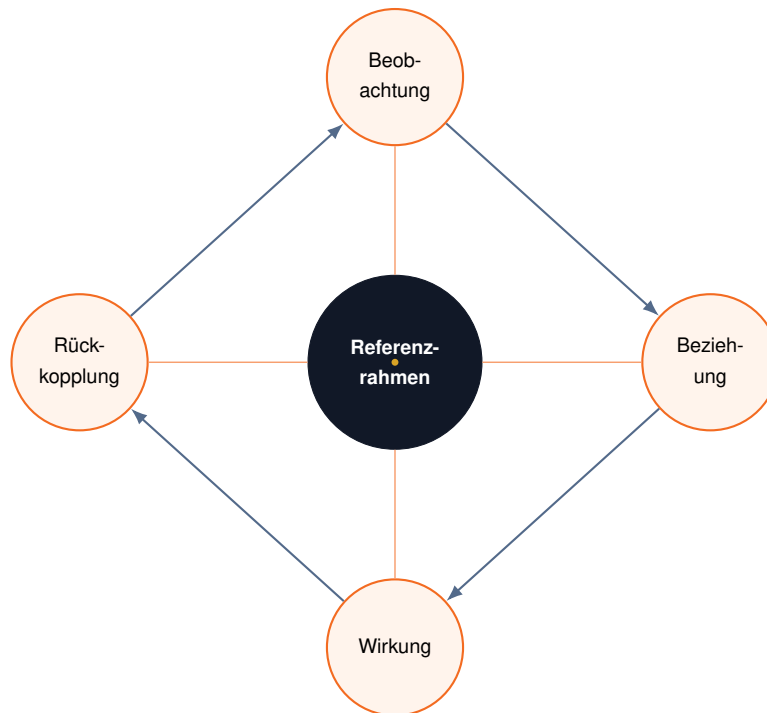


Abbildung 1: ADL als geschlossener Beobachtungs-, Bewertungs- und Rückkopplungskreis.

Der **Erkenntnisvertrag** einer Anwendung besteht aus fünf vor der Auswertung festzuhaltenden Angaben:

1. Welche Größen sind direkt gemessen, welche geschätzt und welche normativ gesetzt?
2. Welche Einheiten, Skalen, Zeitabstände und Unsicherheiten gelten?
3. Bedeutet eine Kante Korrelation, prognostische Kopplung, Mechanismus oder bewusst gesetzte Designbeziehung?
4. Welche Beobachtung würde die gewählte Modellierung schwächen oder widerlegen?
5. Welche reale Entscheidung bleibt außerhalb der automatischen Berechnung?

2. Mathematische Grundobjekte

Für eine Hierarchieebene $\ell \in \mathbb{N}_0$ und diskrete Zeitpunkte $t \in \mathbb{N}_0$ sei

$$\mathcal{G}_t^{(\ell)} = (V^{(\ell)}, E_t^{(\ell)}, A_t^{(\ell)}) \quad (1)$$

ein gewichteter gerichteter Graph mit $n_\ell = |V^{(\ell)}|$ Elementen. Sofern keine Verwechslung droht, wird der Ebenenindex unterdrückt.

2.1 Systemzustand und Beobachtung

Der latente Systemzustand sei $x_t^\circ \in \mathcal{X} \subseteq \mathbb{R}^n$. Beobachtbar ist im Allgemeinen nur

$$y_t = h_t(x_t^\circ) + \varepsilon_t, \quad (2)$$

wobei h_t die Messabbildung und ε_t Messfehler oder nicht modellierte Einflüsse bezeichnet. Aus $y_{0:t}$ entsteht ein Schätzwert $\hat{x}_t$ mit Unsicherheitsangabe, etwa einer Kovarianzmatrix $\Sigma_{x,t}$. In den nachfolgenden Kerngleichungen steht x_t für das in der konkreten Anwendung ausdrücklich deklarierte Datenprodukt: direkte Messung, Schätzung oder codierter Zustand.

2.2 Beziehungsmatrix

$A_t = [a_{ij,t}] \in \mathbb{R}^{n \times n}$ beschreibt die modellierte gerichtete Kopplung. Es gilt die Konvention

$$a_{ij,t} : j \longrightarrow i. \quad (3)$$

Das Vorzeichen bezeichnet die modellierte Wirkungsrichtung, der Betrag die Stärke innerhalb der gewählten Skalierung. Selbstkanten $a_{ii,t}$ sind zulässig, müssen aber begründet werden.

Ontische und epistemische Änderung

Wenn A_t eine geschätzte Struktur ist, bedeutet $A_{t+1} \neq A_t$ zunächst eine Änderung unseres Modells, nicht zwingend eine reale Änderung des Systems. Eine Anwendung muss kennzeichnen, ob A_t einen angenommenen Mechanismus, einen statistischen Prädiktor oder eine Designvorgabe repräsentiert.

2.3 Bewertungsrahmen

Der Bewertungsrahmen wird getrennt vom System geführt:

$$\Theta_t = (z_t, w_t^*, G_t, D_{x,t}, D_{w,t}, D_{g,t}, Q_t, R_t, T_t, \alpha_t, \beta_t, \gamma_t, \delta_t). \quad (4)$$

Dabei ist $z_t \in \mathbb{R}^n$ der Referenzzustand, $w_t^* \in \mathbb{R}^n$ die Referenzwirkung und

$$G_t = \{x \in \mathcal{X} : g_{k,t}(x) \leq 0, k = 1, \dots, K\} \quad (5)$$

die zulässige Zustandsmenge. Die positiven Diagonalmatrizen D_x, D_w, D_g liefern Bezugsskalen; $Q, R, T \succeq 0$ gewichten Richtungen und Wechselwirkungen; $\alpha, \beta, \gamma, \delta \geq 0$ gewichten die vier Bewertungsblöcke.

Warum der Rahmen ein eigenes Objekt ist

Ein System kann unverändert bleiben, während Ziel, Grenze, Gewicht oder Messskala verändert wird. Würden diese Größen im Zustand versteckt, sähe eine geänderte Bewertung wie eine reale Systembewegung aus. Die Trennung q_t versus Θ_t macht diese Verschiebung sichtbar.

3. Wirkung auf dem gerichteten Graphen

Die allgemeine relationale Wirkungsabbildung lautet

$$w_t = \Psi(x_t, A_t; b_t), \quad (6)$$

$$w_{i,t} = b_{i,t} + \sum_{j=1}^n a_{ij,t} \phi_{ij,t}(x_{j,t}, x_{i,t}). \quad (7)$$

b_t erfasst modellierte externe oder autonome Beiträge. $\phi_{ij,t}$ erlaubt nichtlineare, sättigende, schwellenabhängige oder zustandsabhängige Wechselwirkungen. Der lineare Spezialfall ist

$$\phi_{ij,t}(x_j, x_i) = x_j, \quad w_t = A_t x_t + b_t. \quad (8)$$

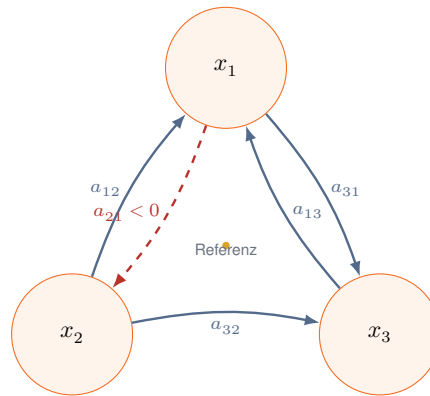


Abbildung 2: Gewichtete gerichtete Beziehungen; die Matrix wird zeilenweise als eingehende Wirkung gelesen.

3.1 Interpretationsgrenze

Die Gleichung berechnet eine Wirkung unter der gewählten Struktur. Sie beweist nicht, dass die Kanten kausal sind. Für eine kausale Lesart sind zusätzlich ein Identifikationsdesign, plausible Ausschlussannahmen, Interventionsdaten oder ein begründetes Strukturmodell nötig. Ohne diese Voraussetzungen ist A_t als beschreibende oder prognostische Kopplung zu lesen.

4. Dimensionslose Abweichungsfunktion

Damit Temperatur, Geld, Wahrscheinlichkeit oder Belastung nicht allein aufgrund ihrer Einheit gegeneinander gewichtet werden, werden alle Residuen zunächst skaliert:

$$r_{x,t} = D_{x,t}^{-1}(x_t - z_t), \quad (9)$$

$$r_{w,t} = D_{w,t}^{-1}(w_t - w_t^*), \quad (10)$$

$$r_{\Delta,t} = D_{x,t}^{-1}(x_t - x_{t-1}), \quad (11)$$

$$\tilde{g}_{k,t}(x) = g_{k,t}(x)/d_{g,k,t}. \quad (12)$$

Hierbei müssen alle Diagonaleinträge der Skalenmatrizen strikt positiv sein. Die dimensionslose Grenzstrafe lautet

$$P_{G_t}(x) = \sum_{k=1}^K [\tilde{g}_{k,t}(x)]_+^2, \quad [a]_+ = \max(0, a). \quad (13)$$

Für den auswertbaren Zustand

$$q_t = (x_t, x_{t-1}, A_t) \quad (14)$$

definiert ADL das Abweichungsfunktional

$$J(q_t; \Theta_t) = \alpha_t r_{x,t}^\top Q_t r_{x,t} + \beta_t r_{w,t}^\top R_t r_{w,t} + \gamma_t P_{G_t}(x_t) + \delta_t r_{\Delta,t}^\top T_t r_{\Delta,t}. \quad (15)$$

Block	Frage	Fehlinterpretation, die vermieden wird
Zustandsabweichung	Wie weit liegt der Zustand vom deklarierten Referenzzustand?	Nähe zum Zentrum ist nicht automatisch sachliche Güte.
Wirkungsabweichung	Welche Folgen erzeugt die Beziehungsstruktur relativ zur Referenzwirkung?	Ähnliche Zustände können verschiedene Wirkungen haben.
Grenzverletzung	Welche expliziten Nebenbedingungen werden überschritten?	Eine Strafe ersetzt keine rechtliche oder ethische Begründung.
Instabilität	Wie stark ändert sich der Zustand zwischen zwei Messpunkten?	Veränderung ist nicht per se schlecht; der Term misst nur Ausmaß.

Tabelle 1: Die vier unterscheidbaren Beiträge von J .

4.1 Relativer Balancewert

$$B(q_t; \Theta_t) = \frac{1}{1 + J(q_t; \Theta_t)}. \quad (16)$$

Proposition 1 – Wertebereich

Sind $Q_t, R_t, T_t \succeq 0$, alle Blockgewichte nichtnegativ und alle Skalen strikt positiv, dann gilt $J \geq 0$ und damit $0 < B \leq 1$. Der Wert $B = 1$ bedeutet lediglich, dass sämtliche positiv gewichteten Residuen im Modell verschwinden. Sind Gewichtsmatrizen singulär oder Blockgewichte null, kann $J = 0$ trotz Abweichungen in unbewerteten Richtungen auftreten.

Der Wert B ist weder Wahrscheinlichkeit noch universeller Balancegrad. Er ist eine monotone, leichter lesbare Darstellung derselben Information wie J und bleibt an Θ_t gebunden.

5. Dynamik, Rückkopplung und Lernen

Eine allgemeine diskrete Zustandsdynamik ist

$$x_{t+1} = \Pi_{G_{t+1}}(f_{\vartheta}(x_t, A_t, e_t) + M_t u_t + \xi_t). \quad (17)$$

e_t bezeichnet exogene Eingaben, u_t den tatsächlich autorisierten Eingriff, M_t dessen Abbildung in Zustandskoordinaten und ξ_t nicht modellierte Dynamik. Eine Euler-Diskretisierung ist der Spezialfall $f_{\vartheta}(x, A, e) = x + \Delta t F_{\vartheta}(x, A, e)$.

Die Projektion ist definiert durch

$$\Pi_G(v) \in \arg \min_{y \in G} \|y - v\|_2. \quad (18)$$

Ist G nichtleer, abgeschlossen und konvex, ist die euklidische Projektion eindeutig. Bei nichtkonvexen Mengen kann sie mehrwertig sein; dann muss eine Auswahlregel offengelegt werden.

5.1 Update der modellierten Beziehungen

Die Beziehungsstruktur wird innerhalb einer zulässigen Matrixmenge $\mathcal{A}_t$ aktualisiert:

$$A_{t+1} = \Pi_{\mathcal{A}_{t+1}}(A_t + \mu_t H_{\vartheta}(\mathcal{J}_t)), \quad (19)$$

wobei $\mathcal{J}_t$ die bis t verfügbare Information und $\mu_t \geq 0$ die Lernrate ist. $\mathcal{A}_t$ kann Vorzeichen, Sparsität, Maximalstärken oder verbotene Kanten festlegen.

Eine mögliche, aber nicht zwingende Spezifikation ist ein projizierter Gradienten-Schritt auf einer Prognoseverlustfunktion:

$$\mathcal{L}_t(A) = \|D_y^{-1}(y_{t+1} - \hat{y}_{t+1}(A))\|_2^2 + \lambda_A \|A\|_1, \quad (20)$$

$$A_{t+1} = \Pi_{\mathcal{A}}(A_t - \mu_t \nabla_A \mathcal{L}_t(A_t)). \quad (21)$$

Die ℓ_1 -Strafe kann eine sparsame Struktur fördern, beweist aber keine kausale Kante.

5.2 Minimalbedingungen für eine eindeutige Rechenvorschrift

- f_{ϑ} , Ψ und H_{ϑ} sind auf den verwendeten Mengen definiert; für Stabilitäts- oder Konvergenzaussagen sind zusätzliche Regularitätsannahmen erforderlich.
- G_t und $\mathcal{A}_t$ sowie jede Auswahlregel bei Mehrdeutigkeit sind dokumentiert.
- Zeitraster, Eingriffseinheit und Abbildung M_t sind dimensionskonsistent.
- Mess-, Prozess- und Parametersicherheit werden nicht als Null behandelt, wenn sie faktisch unbekannt sind.

6. Drift ohne verschobene Torpfosten

Die einfache Differenz $J_{t+1} - J_t$ ist nur dann eindeutig interpretierbar, wenn Beziehungsschätzung und Bewertungsrahmen unverändert bleiben. ADL 2.0 führt deshalb eine exakte dreiteilige Driftbuchhaltung ein.

Definiere die Zwischenzustände

$$q_{t+1}^x = (x_{t+1}, x_t, A_t), \quad (22)$$

$$q_{t+1}^A = (x_{t+1}, x_t, A_{t+1}). \quad (23)$$

Dann gelten

$$D_t^x = J(q_{t+1}^x; \Theta_t) - J(q_t; \Theta_t), \quad \text{Systembewegung bei eingefrorenem Modell und Rahmen,} \quad (24)$$

$$D_t^A = J(q_{t+1}^A; \Theta_t) - J(q_{t+1}^x; \Theta_t), \quad \text{Beziehungs- oder Modellupdate,} \quad (25)$$

$$D_t^\Theta = J(q_{t+1}^A; \Theta_{t+1}) - J(q_{t+1}^A; \Theta_t), \quad \text{Referenzrahmenverschiebung.} \quad (26)$$

$$D_t^{\text{tot}} = J(q_{t+1}^A; \Theta_{t+1}) - J(q_t; \Theta_t) = D_t^x + D_t^A + D_t^\Theta. \quad (27)$$

Proposition 2 – Exakte Zerlegung

Die Identität folgt durch zweimaliges Addieren und Subtrahieren desselben Zwischenwertes. Sie benötigt keine probabilistische oder dynamische Zusatzannahme. Ihre Aussage ist buchhalterisch: Jede gemessene Gesamtänderung wird der gewählten Reihenfolge nach genau auf drei Veränderungsquellen verteilt.

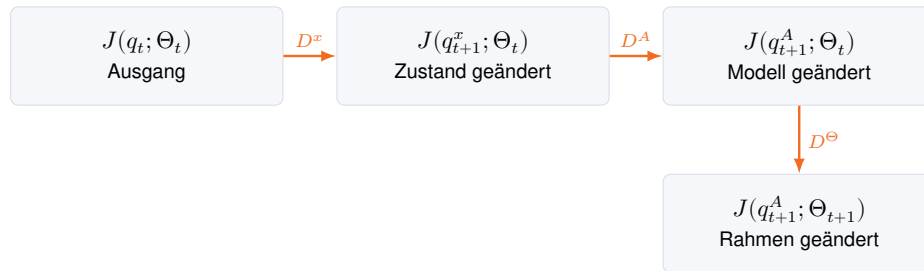


Abbildung 3: Die Driftzerlegung verhindert, dass Systembewegung, Modelllernen und geänderte Zielsetzung verwechselt werden.

Ein negativer Wert von D^x bezeichnet eine Annäherung unter dem alten Rahmen; ein positiver Wert eine Entfernung. Ein negatives D^Θ kann hingegen allein dadurch entstehen, dass der Maßstab nachträglich an den beobachteten Zustand angepasst wurde. Das ist keine Systemverbesserung.

7. Menschlich kontrollierte Intervention

Die formale Trennung zwischen Berechnung und Entscheidung ist Bestandteil des Modells. Order A erzeugt aus der Informationssmenge $\mathcal{J}_t$ transparente Kandidaten, Prognosen und Unsicherheiten. Order M bezeichnet die menschliche Instanz, die auswählt, verändert, vertagt oder verwirft.

Für $u \in \mathcal{U}_t$ kann ein risikosensitiver Kandidatenwert berechnet werden:

$$\widehat{L}_t(u) = \mathbb{E}[J(\tilde{q}_{t+1}(u); \Theta_t) \mid \mathcal{J}_t] + \lambda_c c_t(u) + \kappa \text{SD}[J(\tilde{q}_{t+1}(u); \Theta_t) \mid \mathcal{J}_t]. \quad (28)$$

$c_t(u)$ bildet offengelegte Interventionskosten ab; κ bestimmt die Gewichtung prognostischer Unsicherheit. Statt einen einzigen Skalar zu erzwingen, kann Order A die Pareto-Menge aus erwarteter Abweichung, Kosten und Risiko anzeigen.

$$\mathcal{U}_t^A = \text{Suggest}_A(\mathcal{U}_t, \widehat{L}_t, \mathcal{J}_t), \quad (29)$$

$$u_t = \text{Choose}_M(\mathcal{U}_t^A \cup \{0\}, \mathcal{J}_t). \quad (30)$$

Rollenformel

ADL stellt die Syntax und Prüfgrößen bereit. Order A simuliert, unterscheidet und schlägt vor. Order M bezeichnet den gültigen Kontext und verantwortet die reale Auswahl.

Die menschliche Entscheidung ist kein automatischer Wahrheitsgarant. Automation Bias, Zeitdruck, Machtasymmetrien oder fehlende Fachkenntnis können auch eine menschlich bestätigte Intervention verschlechtern. Deshalb gehören zum Protokoll: sichtbare Annahmen, Alternativen einschließlich Nichtstun, Unsicherheit, Begründung, Freigabestatus und nachträgliche Prüfung durch D^x .

Phase	Order A	Order M
Bezeichnung	erkennt Datenlücken und schlägt Variablen vor	bestätigt Bedeutung, Systemgrenze und Verantwortung
Simulation	berechnet Kandidaten, Baselines und Unsicherheit	bestimmt zulässige Alternativen und Kosten
Auswahl	zeigt Rangfolge oder Pareto-Menge	wählt, ändert, verwirft oder vertagt
Rückkopplung	aktualisiert Prognosefehler und Modellhinweise	beurteilt reale Nebenfolgen und stoppt bei Bedarf

Tabelle 2: Funktionsgrenze zwischen Assistenz und Entscheidung.

8. Rekursion und Mehrskaligkeit

Eine Teilordnung darf nicht allein deshalb zum Makroelement erklärt werden, weil ihre Elemente räumlich nah oder sprachlich ähnlich sind. Die Verdichtung benötigt eine explizite Abbildung und einen gemessenen Informationsverlust.

Sei $C^{(\ell)} \in \mathbb{R}^{m \times n}$ eine nichtnegative, zeilenstochastische Aggregationsmatrix und $D^{(\ell)} \in \mathbb{R}^{n \times m}$ eine Rekonstruktionsabbildung mit

$$C^{(\ell)} \mathbf{1}_n = \mathbf{1}_m, \quad C^{(\ell)} D^{(\ell)} = I_m. \quad (31)$$

Im linearen Fall wird definiert:

$$x_t^{(\ell+1)} = C^{(\ell)} x_t^{(\ell)}, \quad (32)$$

$$A_t^{(\ell+1)} = C^{(\ell)} A_t^{(\ell)} D^{(\ell)}. \quad (33)$$

Die Makrodynamik ist nur dann geschlossen, wenn der aggregierte nächste Zustand hinreichend gut aus Makrogrößen vorhergesagt werden kann. Ein relativer Wirkungserhaltungsfehler ist

$$\varepsilon_C(t) = \frac{\|C A_t x_t - A_t^{(\ell+1)} C x_t\|_2}{\|C A_t x_t\|_2 + \epsilon}, \quad \epsilon > 0. \quad (34)$$

Für nichtlineare Ψ wird die Makroabbildung $\Psi^{(\ell+1)}$ separat gelernt und derselbe Grundsatz out-of-sample geprüft.

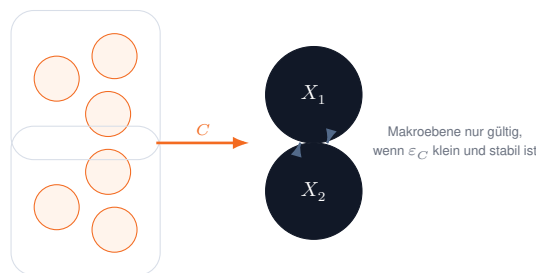


Abbildung 4: Verdichtung mit explizitem Aggregationsoperator und prüfbarem Verlust.

Eine Teilordnung gilt für eine konkrete Anwendung erst dann als stabil verdichtbar, wenn die vorab gesetzten Kriterien über ein Validierungsfenster erfüllt sind, etwa:

- geringe interne Streuung relativ zur gewählten Funktion,
- zeitliche Beständigkeit der Gruppenzugehörigkeit,
- niedriger out-of-sample Aggregationsfehler ε_C ,
- kein Verlust einer für Prognose oder Entscheidung relevanten Minderheitsstruktur.

9. Vollständig gerechnetes Beispiel

Das Beispiel prüft ausschließlich die Rechenlogik. Die drei normierten Elemente heißen Wissen, Energie und Fokus. Es gilt $G = [0, 1]^3$, $D_x = D_w = I$, $T = I$, $b = 0$ und die Matrixkonvention $a_{ij} : j \rightarrow i$.

9.1 Ausgangsdaten und Wirkung

$$x_t = \begin{pmatrix} 0.40 \\ 0.70 \\ 0.30 \end{pmatrix}, \quad x_{t-1} = \begin{pmatrix} 0.35 \\ 0.75 \\ 0.25 \end{pmatrix}, \quad z = \begin{pmatrix} 0.80 \\ 0.60 \\ 0.70 \end{pmatrix}, \quad w^* = \begin{pmatrix} 0.50 \\ 0.10 \\ 0.50 \end{pmatrix}. \quad (35)$$

$$A = \begin{pmatrix} 0 & 0.3 & 0.5 \\ -0.2 & 0 & 0.2 \\ 0.4 & 0.3 & 0 \end{pmatrix}, \quad w_t = Ax_t = \begin{pmatrix} 0.36 \\ -0.02 \\ 0.37 \end{pmatrix}. \quad (36)$$

Die Matrixmultiplikation ist zeilenweise überprüfbar:

$$\begin{aligned} w_1 &= 0(0.4) + 0.3(0.7) + 0.5(0.3) = 0.36, \\ w_2 &= -0.2(0.4) + 0(0.7) + 0.2(0.3) = -0.02, \\ w_3 &= 0.4(0.4) + 0.3(0.7) + 0(0.3) = 0.37. \end{aligned}$$

9.2 Abweichungsfunktion

Verwendet werden

$$Q = \text{diag}(1, 0.5, 2), \quad R = I, \quad (\alpha, \beta, \gamma, \delta) = (1, 2, 5, 0.5). \quad (37)$$

Da $x_t \in G$, ist $P_G(x_t) = 0$. Die vier Beiträge sind

$$J_x = (x_t - z)^T Q (x_t - z) = 0.485000, \quad (38)$$

$$J_w = 2(w_t - w^*)^T (w_t - w^*) = 2(0.050900) = 0.101800, \quad (39)$$

$$J_G = 5P_G(x_t) = 0, \quad (40)$$

$$J_\Delta = 0.5 \|x_t - x_{t-1}\|_2^2 = 0.5(0.007500) = 0.003750. \quad (41)$$

Somit

$$J_t = 0.590550, \quad B_t = \frac{1}{1 + 0.590550} = 0.628713. \quad (42)$$

Reproduzierbarkeitskorrektur gegenüber Version 1.0

Die Zielwirkung $w^* = (0.50, 0.10, 0.50)^T$ war in der ersten Fassung im Rechenabschnitt nicht ausgewiesen. Ohne sie ließ sich der Wirkungsbeitrag 0.1018 nicht vollständig nachrechnen. Version 2.0 schließt diese Lücke und zeigt alle Teilterme.

9.3 Simulierte und menschlich autorisierte Intervention

Für diesen einzelnen Prüfschritt seien $f(x, A, e) = x$, $M = I$, $\xi = 0$, $A_{t+1} = A_t$ und $\Theta_{t+1} = \Theta_t$. Order M wählt aus sichtbaren Kandidaten

$$u_t = \begin{pmatrix} 0.08 \\ -0.02 \\ 0.10 \end{pmatrix}, \quad x_{t+1} = x_t + u_t = \begin{pmatrix} 0.48 \\ 0.68 \\ 0.40 \end{pmatrix}. \quad (43)$$

Der Folgewirkungsvektor lautet

$$w_{t+1} = Ax_{t+1} = \begin{pmatrix} 0.404 \\ -0.016 \\ 0.396 \end{pmatrix}. \quad (44)$$

Damit ergeben sich

$$J_{x,t+1} = 0.285600, \quad (45)$$

$$J_{w,t+1} = 2(0.033488) = 0.066976, \quad (46)$$

$$J_{G,t+1} = 0, \quad (47)$$

$$J_{\Delta,t+1} = 0.5(0.016800) = 0.008400, \quad (48)$$

$$J_{t+1} = 0.360976, \quad B_{t+1} = 0.734767. \quad (49)$$

Da Modell und Rahmen eingefroren sind,

$$D_t^x = J_{t+1} - J_t = -0.229574, \quad D_t^A = D_t^\Theta = 0. \quad (50)$$

Die Intervention reduziert die modellierte Abweichung in diesem Schritt. Sie beweist weder eine reale Lernwirkung noch die Angemessenheit der gewählten Variablen und Gewichte.

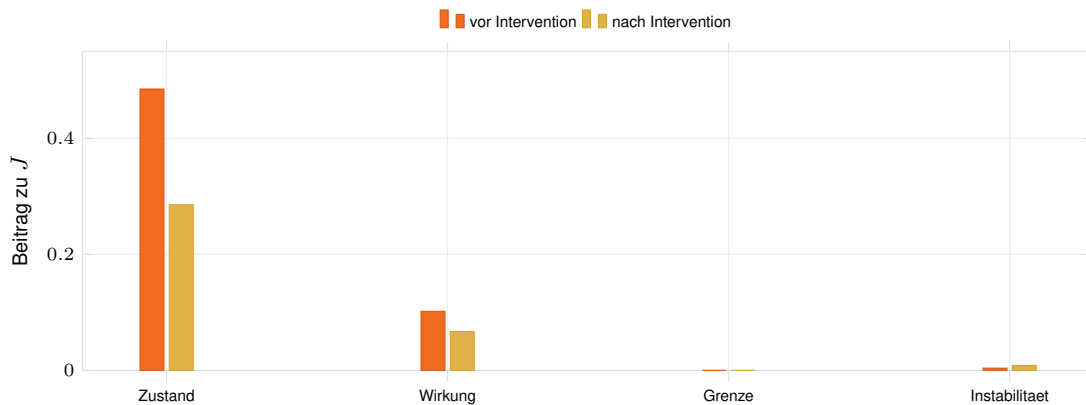


Abbildung 5: Beiträge zu J vor und nach der simulierten Intervention.

10. Prüfung, Falsifizierbarkeit und Anspruchsleiter

Ein formales Modell wird nicht dadurch empirisch belastbar, dass es viele Sachverhalte sprachlich aufnehmen kann. Es muss vor der Beobachtung festlegen, welche Variablen gelten, wie sie gemessen werden, welche Baseline geschlagen werden soll und welches Ergebnis gegen die Modellannahmen spricht.

10.1 Minimaler Prüfplan

1. **Präregistrieren:** Feld, Systemgrenze, Zeittakt, Variablen, Einheiten, h_t , A , Θ , Eingriffe und primäre Gütemaße vor der Testauswertung festlegen.
2. **Trennen:** Daten in Training, Validierung und unangetasteten Test aufteilen; zeitliche Leakage ausschließen.
3. **Vergleichen:** Persistenzmodell, unabhängige Einzelmodelle und mindestens ein alternatives Netzwerk- oder Dynamikmodell als Baselines verwenden.
4. **Ablatieren:** Zustands-, Wirkungs-, Grenz- und Instabilitätsblock einzeln entfernen und den zusätzlichen Nutzen prüfen.
5. **Sensitivität prüfen:** Gewichte, Referenzen, Skalen, Kanten und Messfehler innerhalb begründeter Bereiche variieren.
6. **Out-of-sample testen:** Wirkungsrichtung, Prognosefehler, Interventionsrangfolge, Unsicherheitskalibrierung und Coarse-Graining-Fehler berichten.
7. **Replizieren:** denselben Prüfplan auf neuen Zeiträumen, Gruppen oder Systemen wiederholen.

10.2 Mögliche Widerlegungskriterien

- Die modellierten Kanten sagen Wirkungsrichtung oder Größenordnung wiederholt schlechter als eine einfachere Baseline voraus.

- Interventionen mit niedrigerem prognostiziertem J verschlechtern den vorab definierten realen Zielwert systematisch.
- Die Rangfolge der Interventionen kehrt sich bei kleinen, sachlich plausiblen Änderungen von Θ unkontrolliert um.
- Ein vermeintlicher Fortschritt stammt überwiegend aus $D^\Theta < 0$, während $D^x \geq 0$ bleibt.
- Mehrere deutlich verschiedene Matrizen A erklären die Daten gleich gut; die behauptete Beziehungsstruktur ist nicht identifizierbar.
- Das Coarse-Graining erzeugt hohen out-of-sample Verlust oder verdeckt entscheidungsrelevante Teilstrukturen.
- Ein sparsameres Modell erreicht gleiche oder bessere Prognose und Kalibrierung mit weniger Annahmen.

10.3 Anspruchsleiter

Stufe	Behauptung	Erforderlicher Nachweis	Status hier
0	interne Konsistenz	definierte Größen, dimensionsfähige Formeln, keine Rechenwidersprüche	erfüllt
1	rechnerische Reproduzierbarkeit	Daten und Parameter führen unabhängig zu denselben Zahlen	erfüllt für Beispiel
2	deskriptive Adäquanz	beobachtete Struktur wird mit akzeptablem Fehler abgebildet	offen
3	prognostische Validität	out-of-sample besser und kalibriert gegenüber Baselines	offen
4	Interventionswirksamkeit	prospektive oder kausal begründete Verbesserung	offen
5	Übertragbarkeit	wiederholte Leistung in anderen Zeiten, Gruppen oder Feldern	offen

Tabelle 3: Der Geltungsanspruch wächst nur mit dem jeweiligen Nachweis.

11. Identifizierbarkeit, Unsicherheit und Grenzen

11.1 Identifizierbarkeit

Eine Matrix A ist nicht identifiziert, wenn verschiedene Parameterwerte dieselbe Verteilung der beobachtbaren Daten erzeugen. Dann kann eine präzise aussehende Kante mathematisch unbestimmt sein. Zu prüfen sind insbesondere Kollinearität, zu wenige Zeitpunkte, verdeckte gemeinsame Ursachen, symmetrische Modellalternativen und die Zahl freier Parameter relativ zur Informationsmenge.

11.2 Unsicherheitsrechnung

Mindestens drei Unsicherheitsquellen sind getrennt zu berichten:

- **Messunsicherheit** in y_t und daraus in $\hat{x}_t$,
- **Parameterunsicherheit** in A , f_θ und Ψ ,
- **Rahmenunsicherheit** bei z , w^* , Grenzen und Gewichten.

Mögliche Verfahren sind Bootstrap, Bayes'sche Posteriorverteilungen, Konfidenzintervalle, Ensemblemodelle und globale Sensitivitätsanalyse. Die Methode muss zur Datenstruktur passen; ein Unsicherheitswert ist nur so belastbar wie seine Annahmen.

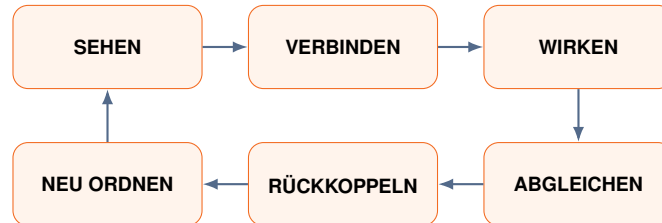
11.3 Grenzen und Nichtbehauptungen

Was aus der Mathematik nicht folgt

- Die Funktionen Ψ , f_θ , H_θ und die Referenzen werden nicht von ADL universell festgelegt.
- Ein hoher Balancewert beweist weder Wahrheit noch Gerechtigkeit, Gesundheit oder moralische Richtigkeit.
- Ein gerichteter Graph beweist keine Kausalstruktur.
- Menschliche Freigabe beseitigt weder Machtfehler noch Automation Bias.
- Rekursion bedeutet nicht, dass jede Skala dieselben Parameter oder Gesetze besitzt.
- Das vorliegende Beispiel zeigt Auswertbarkeit, nicht empirische Wirksamkeit.

Die stärkste wissenschaftliche Funktion der Architektur liegt zunächst in ihrer Prüfbarkeit: Sie zwingt dazu, Beziehungen, Referenzen, Einheiten, Grenzen, Unsicherheiten und Entscheidungszuständigkeit sichtbar zu machen. Gerade diese Sichtbarkeit ermöglicht Kritik, Vergleich und Widerlegung.

12. Rohform für den Kopf



Die rohe Frage

Was wurde tatsächlich beobachtet? Was wirkt im Modell worauf – in welche Richtung und mit welcher Stärke? Gegen welchen offen gelegten Rahmen wird es gemessen? Welche Grenze gilt? Was kommt mit der Zeit zurück? Hat sich das System verändert, das Modell gelernt oder nur der Maßstab verschoben? Wer verantwortet den nächsten Eingriff?

Eine Ordnung wird nicht nach ihrem Namen oder ihrem ersten Erscheinungsbild bewertet, sondern nach der Struktur ihrer Beziehungen, den daraus entstehenden Wirkungen, den geltenden Grenzen und der Richtung, die diese Wirkungen über die Zeit nehmen.

A. Symboltabelle

Symbol	Definition	Interpretationshinweis
t, ℓ	Zeitindex, Hierarchieebene	diskrete Zeit; Mikro-, Meso- oder Makroebene
$\mathcal{G}_t^{(\ell)}$	gewichteter gerichteter Graph	modellierte relationale Ordnung auf Ebene ℓ
x_t°	latenter Zustand	nicht zwingend direkt beobachtbar
y_t	Beobachtung	Messwert nach h_t plus Fehler ε_t
x_t	verwendetes Zustandsprodukt	Messung, Schätzung oder Codierung; muss deklariert sein
$A_t = [a_{ij,t}]$	Beziehungsmatrix	a_{ij} beschreibt $j \rightarrow i$; nicht automatisch kausal
Ψ, ϕ_{ij}	Wirkungsabbildung	linear oder anwendungsbezogen nichtlinear
w_t	modellierter Wirkungsvektor	Ergebnis aus Zustand, Beziehung und externem Beitrag
z_t	Referenzzustand	gesetzter Orientierungspunkt, nicht automatisch Gleichgewicht
w_t^*	Referenzwirkung	deklarierte Ziel- oder Tragfähigkeitswirkung
G_t	zulässige Zustandsmenge	durch $g_{k,t}(x) \leq 0$ beschriebene Grenzen
D_x, D_w, D_g	Skalenmatrizen	machen Residuen dimensionslos; Diagonalen strikt positiv
Q, R, T	Gewichtungsmatrizen	positiv semidefinit; können Richtungen unbewertet lassen
$\alpha, \beta, \gamma, \delta$	Blockgewichte	nichtnegative relative Gewichte der vier J -Blöcke
q_t	auswertbarer Zustand	Tripel (x_t, x_{t-1}, A_t)
Θ_t	Bewertungsrahmen	Referenzen, Grenzen, Skalen und Gewichte
$J(q_t; \Theta_t)$	Abweichungsfunktional	kleiner bedeutet nur unter Θ_t nähere Modellpassung
$B(q_t; \Theta_t)$	relativer Balancewert	$1/(1 + J)$; weder Wahrscheinlichkeit noch absolute Gerechtigkeit
f_ϑ	diskrete Zustandsdynamik	anwendungsbezogen zu spezifizieren und zu validieren
e_t, ξ_t	exogener Eingang, Restdynamik	beobachtete externe Einflüsse und nicht modellierter Anteil
u_t, M_t	Eingriff, Eingriffsabbildung	real autorisierte Handlung und ihre Zustandswirkung
Π_G	Projektion auf G	bei nichtkonvexem G möglicherweise mehrwertig
H_ϑ, μ_t	Update-Regel, Lernrate	verändert die modellierte Beziehungsschätzung
D^x, D^A, D^Θ	Driftkomponenten	System-, Modell- und Rahmenanteil
C, D	Aggregation, Rekonstruktion	bilden Mikro- auf Makroebene und näherungsweise zurück
ε_C	Aggregationsfehler	misst Verlust der Wirkungsdarstellung
$\mathcal{I}_t$	Informationsmenge	Daten, Modell, Annahmen und Unsicherheit bis t
$\mathcal{U}_t^A$	Vorschlagsmenge	von Order A sicht- und prüfbar erzeugte Kandidaten

B. ADL-Kernsystem

$$\begin{aligned}
 \mathcal{G}_t &= (V, E_t, A_t), & w_t &= \Psi(x_t, A_t; b_t), \\
 q_t &= (x_t, x_{t-1}, A_t), & \Theta_t &= (z_t, w_t^*, G_t, D_x, D_w, D_g, Q, R, T, \alpha, \beta, \gamma, \delta), \\
 J(q_t; \Theta_t) &= \alpha r_x^\top Q r_x + \beta r_w^\top R r_w + \gamma P_G(x_t) + \delta r_\Delta^\top T r_\Delta, & B_t &= \frac{1}{1 + J(q_t; \Theta_t)}, \\
 x_{t+1} &= \Pi_{G_{t+1}}(f_\vartheta(x_t, A_t, e_t) + M_t u_t + \xi_t), & A_{t+1} &= \Pi_{\mathcal{A}_{t+1}}(A_t + \mu_t H_\vartheta(J_t)), \\
 \mathcal{U}_t^A &= \text{Suggest}_A(\mathcal{U}_t, \widehat{L}_t, \mathcal{J}_t), & u_t &= \text{Choose}_M(\mathcal{U}_t^A \cup \{0\}, \mathcal{J}_t), \\
 D_t^{\text{tot}} &= D_t^x + D_t^A + D_t^\Theta, & x^{(\ell+1)} &= C^{(\ell)} x^{(\ell)}, \quad A^{(\ell+1)} = C^{(\ell)} A^{(\ell)} D^{(\ell)}.
 \end{aligned}$$

(51)

Lesereihenfolge

Beobachtung beziehungsweise Schätzung erzeugt x_t . Die deklarierte Beziehung A_t erzeugt über Ψ eine Wirkung w_t . Der getrennte Rahmen Θ_t bewertet Zustand, Wirkung, Grenzen und zeitliche Veränderung. Order A simuliert zulässige Kandidaten und Unsicherheit; Order M autorisiert höchstens einen realen Eingriff. Die nächste Beobachtung prüft die Systembewegung getrennt von Modell- und Rahmenänderungen. Erst bei kleinem, stabilem Aggregationsfehler wird eine Teilordnung zum Element der nächsthöheren Ebene verdichtet.

Spezifikationspflicht

Das Kernsystem ist eine Modellschablone. Vor jeder Anwendung müssen h_t , Ψ , f_ϑ , H_ϑ , C , sämtliche Variablen, Einheiten, Skalen, Nebenbedingungen, Referenzen, Unsicherheitsmodelle und Update-Regeln operationalisiert werden. Ohne diese Angaben ist keine empirische Aussage bestimmt.

C. Reproduzierbarkeitsprotokoll

Pflichtangabe	Zu archivieren
☐ Systemgrenze und Zweck	begründete Ein- und Ausschlüsse, verantwortliche Stelle
☐ Variablenwörterbuch	Bedeutung, Einheit, Skala, Messzeitpunkt, Transformation
☐ Datenherkunft	Quelle, Auswahlprozess, fehlende Werte, Vorverarbeitung, Version
☐ Graphsemantik	Bedeutung jeder Kantenart, Richtungskonvention, verbotene Kanten
☐ Bewertungsrahmen	$z, w^*, G, D_x, D_w, D_g, Q, R, T$ und Blockgewichte
☐ Unsicherheit	Mess-, Parameter- und Rahmenunsicherheit samt Verfahren
☐ Dynamik	f_{ϑ} , Zeittakt, Eingriffe, M , Projektion, Restterm
☐ Lernregel	H_{ϑ} , Lernrate, Regularisierung, zulässige Matrixmenge
☐ Baselines	mindestens Persistenz und ein sachlich alternatives Modell
☐ Prüfdesign	Train/Validierung/Test, Metriken, Präregistrierung, Abbruchkriterien
☐ Entscheidung	Kandidaten, Nichtstun-Option, Kosten, Unsicherheit, menschliche Begründung
☐ Rückkopplung	D^x, D^A, D^{Θ} , Nebenfolgen, Modellfehler, Korrekturhistorie
☐ Mehrskaligkeit	C, D , Stabilitätskriterien, ε_C in und out-of-sample
☐ Rechenartefakte	Code, Zufallsseed, Bibliotheksversionen, Roh- und Ergebnisdaten

Numerischer Rechencheck der Referenzinstanz

Prüfgröße	vorher	nachher	Rechenregel
Zustandsblock	0.485000	0.285600	$r_x^T Q r_x$
Wirkungsblock	0.101800	0.066976	$2 r_w^T r_w$
Grenzblock	0	0	$5 P_G(x)$
Instabilitätsblock	0.003750	0.008400	$0.5 r_{\Delta}^T r_{\Delta}$
Gesamtwert J	0.590550	0.360976	Summe der vier Blöcke
Balancewert B	0.628713	0.734767	$1/(1 + J)$

Akzeptanzbedingung für eine unabhängige Implementierung

Mit den in Abschnitt 9 angegebenen Vektoren und Matrizen müssen die ausgewiesenen Blockwerte bis auf Rundungsfehler reproduziert werden. Außerdem muss $0.360976 - 0.590550 = -0.229574$ gelten. Abweichungen sind vor jeder empirischen Verwendung als Implementierungs- oder Spezifikationsfehler zu klären.

D. Literatur und methodische Anschlüsse

Literatur

- [1] Newman, M. E. J. (2018). *Networks*. 2. Auflage. Oxford University Press. DOI: [10.1093/oso/9780198805090.001.0001](https://doi.org/10.1093/oso/9780198805090.001.0001).
- [2] Ljung, L. (1999). *System Identification: Theory for the User*. 2. Auflage. Prentice Hall.
- [3] Khalil, H. K. (2002). *Nonlinear Systems*. 3. Auflage. Prentice Hall.
- [4] Bertsekas, D. P. (2016). *Nonlinear Programming*. 3. Auflage. Athena Scientific.
- [5] Gorban, A. N., Kazantzi, N., Kevrekidis, I. G., Öttinger, H.-C. und Theodoropoulos, C. (Hrsg.) (2006). *Model Reduction and Coarse-Graining Approaches for Multiscale Phenomena*. Springer. DOI: [10.1007/3-540-35888-9](https://doi.org/10.1007/3-540-35888-9).
- [6] Wu, X. et al. (2022). A survey of human-in-the-loop for machine learning. *Future Generation Computer Systems*, 135, 364–381. DOI: [10.1016/j.future.2022.05.014](https://doi.org/10.1016/j.future.2022.05.014).
- [7] Saltelli, A., Ratto, M., Andres, T., Campolongo, F., Cariboni, J., Gatelli, D., Saisana, M. und Tarantola, S. (2008). *Global Sensitivity Analysis: The Primer*. Wiley. DOI: [10.1002/9780470725184](https://doi.org/10.1002/9780470725184).
- [8] Pearl, J. (2009). *Causality: Models, Reasoning, and Inference*. 2. Auflage. Cambridge University Press.
- [9] Box, G. E. P. (1976). Science and Statistics. *Journal of the American Statistical Association*, 71(356), 791–799. DOI: [10.1080/01621459.1976.10480949](https://doi.org/10.1080/01621459.1976.10480949).

E. Autor und Projekt

M. Selman Yildizhan entwickelt ADL – Architektur des Lebens als relationale Syntax zur Darstellung von Elementen, Beziehungen, Wirkungen, Referenzrahmen, Grenzen, Rückkopplungen und rekursiver Ordnungsbildung. Order A bezeichnet die assistierende Analyse- und Vorschlagsfunktion; Order M die menschliche Instanz der Bezeichnung, Freigabe und Entscheidung. Das Projekt wird unter Epoch Creation in Hamburg entwickelt.

Kontakt und Version

Epoch Creation | Hamburg, Germany

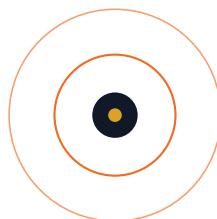
epochcreationservice@gmail.com

Version 2.0 | 14. Juli 2026

Projekt-DOI: [10.5281/zenodo.21209954](https://doi.org/10.5281/zenodo.21209954)

Publikationshinweis

Diese Fassung ersetzt inhaltlich Version 1.0. Änderungen: vollständige Angabe der Zielwirkung im Rechenbeispiel, dimensionslose Residuen, Trennung von Beobachtung und latentem Zustand, projizierte Beziehungsupdates, exakte Driftzerlegung, risikosensitive Interventionsvorschläge, prüfbares Coarse-Graining, Anspruchsleiter, Reproduzierbarkeitsprotokoll und Literaturanschlüsse.



Ende der mathematischen Grundlegung – Version 2.0