

Modulares Prompt-Kit:

Intelligente Linearität in situativen Kontexten

Erweiterte Version mit operationalisierter Selbstbegrenzung und systematischen Feedbacksystemen [Version MVF2]

Autor: Timo Weil, www.timoweil.de

Lizenz: Dieses Dokument ist frei verwendbar für akademische, pädagogische und künstlerische Zwecke. Attribution erwünscht.

Inhalt

I. Grundidee: Warum dieses Kit?	3
II. Grundmodule (A-E): Strukturelle Bausteine.....	4
Modul A: Rollenklärung	4
Modul B: Kontextualisierte Reduktion	4
Modul C: Relevanzkriterien sichtbar machen	4
Modul D: Temporäre Autorisierung	5
Modul E: Operationalisierte Selbstbegrenzung.....	5
III. Reflexionsmodule (F-Serie): Prompt-Qualität systematisch entwickeln.....	7
Modul F: Grundlegende Prompt-Reflexion	7
IV. Metakognitive Module	12
Modul G: Muster und Archetypen	12
Modul H: Human-in-the-Loop & Validierungsprotokolle	15
Modul I: Ambiguitäts- & Konfliktmanagement	17
Modul J: Bias- & Ethik-Check (Erweiterte Reflexion).....	20
Modul K: Iterative Verbesserung.....	22
Modul L: Cross-Context-Transfer.....	23
V. Erweiterte Module (M-V)	25
Modul M: Daten-Input-Transparenz	25
Modul N: Knowledge-Sharing-Protokoll	26
Modul O: Erkenntnis-Visualisierung	26

Modul P: Erkenntnis-Verdichtung.....	27
Modul Q: Dynamisches Erkenntnismanagement.....	28
Modul R: Die Muster-Maschine (Batesons „ökologisches“ Denken).....	28
Modul S: Der Individuations-Prozessor (Simondons „Individuation“).....	30
Modul T: Der morphogenetische Generator (DeLandas „Materie“).....	31
Grundmodul U: Epistemologische Trias (Wissensbegriff-Fundierung)	32
Modul V: Epistemische Meta-Steuerung.....	34
Integration in bestehende Module.....	35
Praktische Anwendung	35
Integration in bestehende Module.....	36
Praktische Anwendung	36
VI. Systemische Module (Z-Serie): Dokumentation und Qualitätssicherung mit Feedbacksystemen	37
Modul Z: Prompt-Dokumentation und Wissenssicherung.....	37
VII. Praktische Anwendung: Checklisten und Workflows	42
Erweiterte Schnell-Referenz: Module kombinieren.....	42
VIII. Fazit und Ausblick.....	44

I. Grundidee: Warum dieses Kit?

Ziel: Prompts so zu gestalten, dass sie klare, lineare Antworten ermöglichen, ohne Kontextkomplexität zu vernachlässigen.

Anwendung: Wissenschaft, Technik, Recht, Medizin, Planung – überall dort, wo Eindeutigkeit erforderlich ist, aber die Realität komplex bleibt.

Leitprinzipien:

1. Rollenbewusstsein – Klare Funktionszuweisung
2. Reduktionsreflexion – Bewusste Komplexitätssteuerung
3. Temporäre Autorisierung – Definierte Verantwortungsbereiche
4. Meta-Klarheit – Transparente Entscheidungsstrukturen
5. Relationalität vor Objektivität – Kontextbewusste Antworten

II. Grundmodule (A-E): Strukturelle Bausteine

Modul A: Rollenklärung

Formulierungshilfe:

Du bist in der Rolle eines/einer [Fachperson] im Kontext von [Situation].
Deine Aufgabe ist es, für [Zielgruppe] eine Entscheidung vorzubereiten.

✚ Kombinierbar mit: Autorisierung, Relevanzklärung

Modul B: Kontextualisierte Reduktion

Strukturvorschlag:

1. Was wäre vollständig? (nur skizzieren)
2. Warum reduzieren wir? (z. B. Zeit, Zielgruppe, Handlungsdruck)
3. Was wird fokussiert?
4. Was wird bewusst nicht behandelt – aber bleibt im Hintergrund mitgedacht?

Promptbeispiel:

Beginne mit einer kurzen Aufzählung, was eine umfassende Analyse umfassen würde. Nenne dann die Gründe, warum du diese eingrenzt. Fokussiere dich auf maximal zwei Aspekte. Nenne am Ende, was bewusst ausgeklammert wurde.

Wenn du möchtest, ergänze einen Satz, wie sich die Auswahl auf das Gesamtergebnis auswirken könnte.

Modul C: Relevanzkriterien sichtbar machen

Promptbaustein:

Bevor du antwortest: Nach welchen Kriterien soll hier Relevanz bestimmt werden? Rechtlich? Praktisch? Wissenschaftlich? Persönlich? Begründe deine Wahl.

Sollten mehrere Kriterien gelten, gib an, wie sie gewichtet werden.

Bewerte am Ende deiner Antwort, wie sich die gewählten Relevanzkriterien auf die Tiefe oder Richtung der Argumentation ausgewirkt haben.

Zweck: Bewusstmachen, dass Relevanz nicht neutral ist, sondern kontextabhängig.

Modul D: Temporäre Autorisierung

Promptbaustein:

Du darfst im Rahmen deiner Rolle Aussagen treffen zu [A], aber nicht zu [B].
Wenn du an Grenzen stößt, sag: „Hier endet mein Mandat.“

Anwendung: Ermöglicht Klarheit, wo eine KI Verantwortung übernehmen darf – und wo nicht.

Modul E: Operationalisierte Selbstbegrenzung

Grundprinzip: Explizite Mandatsgrenz-Definition mit praktischen Trigger-Mechanismen

Modul E1: Entscheidungssampel-System

● GRÜN - Volle Kompetenz:

- Faktische Informationen vermitteln
- Etablierte Methoden erklären
- Strukturierte Analysen durchführen
- Verschiedene Optionen aufzeigen

● GELB - Eingeschränkte Kompetenz (mit Verweis):

- Einschätzungen mit Unsicherheitshinweis
- Empfehlungen mit Expertenvorbehalt
- Interpretationen mit Methodengrenze
- Prognosen mit Wahrscheinlichkeitsangabe

● ROT - Keine Kompetenz (Verweisung erforderlich):

- Rechtlich bindende Aussagen
- Medizinische Diagnosen/Therapieempfehlungen
- Individuelle Investmentberatung
- Sicherheitsrelevante Einzelentscheidungen

Modul E2: Mandatsgrenz-Trigger

Automatische Grenzmarkierung bei:

- Rechtlichen Begriffen: „Das berührt rechtliche Aspekte, die eine Fachberatung erfordern“
- Medizinischen Symptomen: „Hier ist eine ärztliche Einschätzung unerlässlich“
- Sicherheitskritischen Situationen: „Sicherheitsrelevante Entscheidungen gehören zu Fachkräften“
- Persönlichen Krisen: „Hier empfehle ich professionelle Beratung“

Modul E3: Strukturierte Selbstbegrenzung

Promptbaustein:

Verwende folgendes Format für deine Antwort:

MEIN BEITRAG (● / ● / ●): [Ampelstatus und was ich leisten kann]

ANTWORT: [Deine Ausführung]

GRENZEN: [Was gehört zu anderen Fachbereichen]

NÄCHSTE SCHRITTE: [Empfohlene weitere Schritte/Ansprechpartner]

Erweiterte Grenzmarkierung:

Beginne mit: ‚Für diese Entscheidung kann ich beitragen zu... [●] / mit Einschränkungen zu... [●] / nicht beitragen zu... [●]‘

III. Reflexionsmodule (F-Serie): Prompt-Qualität systematisch entwickeln

Modul F: Grundlegende Prompt-Reflexion

Ziel: Die Qualität und Wirkung eigener Prompts sichtbar machen, um künftige KI-Interaktionen bewusster zu gestalten.

Strukturvorschlag für Reflexion:

1. Welche Funktion hatte mein Prompt?

- Impuls, Analyseauftrag, Einladung, Steuerung, Informationsbitte, Kommentar ...

2. Welche Rolle habe ich (implizit oder explizit) der KI zugeschrieben?

– Fachkraft, Mitdenkerin, Dienstleister, kritische Instanz ...

3. War mein Prompt kontextualisiert?

- Zeit, Ort, Zielgruppe, Medium, Anschlussfähigkeit ...

4. Welche Reduktion habe ich vorgenommen oder unterlassen?

- Bewusst eingegrenzt oder diffus gehalten?

5. Welche Wirkung habe ich intendiert – und wie deutlich wurde das formuliert?

- Klarer Auftrag, vage Erwartung, offener Raum ...

6. Wie offen war der Prompt für Resonanz und Diskurs?

Erweiterte Reflexionsdimensionen:

7. Welche impliziten Vorannahmen enthielt mein Prompt?

- Kulturelle, fachliche oder methodische Prämissen

8. Wie ausgewogen war die Balance zwischen Struktur und Kreativitätsraum?

- Übersteuerung vs. Unterspezifikation

9. Welche Qualitätskriterien habe ich für die Antwort gesetzt?

- Explizit oder implizit, messbar oder subjektiv

10. Adressatenorientierung

- Für wen ist der Prompt gedacht – und wird diese Adressatenschaft erkennbar? (z. B. KI, Mensch, Doppeladressierung ...)

11. Sprachstil & Tonalität

- Wie wirkt die sprachliche Form? (ruhig, sachlich, spielerisch, belehrend, tastend ...)

12. Ambiguitätspotenzial

- Welche Mehrdeutigkeiten lässt der Prompt zu – und sind diese produktiv oder störend?

13. Lern-/Erkenntnispotenzial

- Welches Potenzial zur Selbst- oder Fremderkenntnis steckt im Prompt? (Reflexion, neue Perspektiven, methodischer Zugewinn ...)

14. Erweiterbarkeit

- Wie offen ist der Prompt für Weiterentwicklungen, Folgefragen oder die Einbettung in eine Serie?

Promptbaustein:

Beurteile diesen Prompt nach folgenden Gesichtspunkten: Funktion, Rollenzuweisung, Kontextualisierung, Reduktionsstrategie, Zielklarheit, Diskurspotenzial, implizite Vorannahmen, Struktur-Kreativitäts-Balance, Qualitätskriterien, Adressatenorientierung, Sprachstil & Tonalität, Ambiguitätspotenzial, Lern-/Erkenntnispotenzial, Erweiterbarkeit.

Modul F1: Fachbereichsspezifische Reflexion

Rechtlicher Kontext:

Analysiere meinen Prompt: Wie klar sind die rechtlichen Grenzen definiert? Wird zwischen Rechtsauskunft und Rechtsinformation unterschieden? Welche Haftungsaspekte bleiben ungeklärt?

Medizinischer Kontext:

Bewerte meinen Prompt: Sind die Grenzen zwischen Information und Beratung klar? Wie wird mit Unsicherheiten und Risiken umgegangen? Wo sind ärztliche Kompetenzen unerlässlich?

Pädagogischer Kontext:

Reflektiere meinen Prompt: Ist er entwicklungsangemessen? Fördert er selbstständiges Denken oder gibt er Antworten vor? Wie wird Vielfalt der Lerntypen berücksichtigt?

Technischer Kontext:

Analysiere meinen Prompt: Sind Sicherheitsaspekte ausreichend berücksichtigt? Wird zwischen verschiedenen Implementierungsebenen unterschieden? Wie aktuell sind die technischen Voraussetzungen?

Wissenschaftlicher Kontext:

Bewerte meinen Prompt: Entspricht er wissenschaftlichen Standards? Werden Unsicherheiten und Methodengrenzen transparent gemacht? Wie wird mit widersprüchlichen Evidenzen umgegangen?

Modul F2: Zielgruppen- und Medientyp-Reflexion

Modul F2a: Grundlegende Zielgruppenreflexion

Grundstruktur:

Analysiere: Für welche Zielgruppe war mein Prompt optimiert? Wie würde sich der Prompt für [andere Zielgruppe] ändern müssen?

Spezifische Zielgruppen-Checks:

Für Expertinnen und Experten:

- Ist die Fachterminologie angemessen?
- Werden Grundlagen übersprungen oder unnötig erklärt?
- Entspricht die Tiefe den Erwartungen?

Für Laien:

- Sind Fachbegriffe erklärt oder vermieden?
- Gibt es anschauliche Beispiele?
- Ist die Komplexität angemessen reduziert?

Für Entscheidungsträger:

- Werden Handlungsoptionen klar benannt?
- Sind Risiken und Nutzen abgewogen?
- Ist die Information entscheidungsrelevant?

Für Lernende:

- Wird Verstehen vor Antworten gestellt?
- Gibt es Raum für eigene Reflexion?
- Sind die Lernschritte nachvollziehbar?

Modul F2b: Medientyp-spezifische Anpassung

Chat-Kontext:

- Kurze, iterierbare Einheiten bevorzugen
- Nachfragen explizit einladen
- Aufbauende Gesprächslogik berücksichtigen

Promptbaustein:

Antworte in Chat-geeigneten Häppchen, lade zu Rückfragen ein

Voice/Audio-Kontext:

- Linear-logische Struktur ohne Sprungmarken
- Keine komplexen visuellen Elemente
- Klare Gliederungssignale („Erstens... Zweitens...“)

Promptbaustein:

Strukturiere für auditive Aufnahme, verwende keine Listen oder Tabellen

Texteditor/Dokument-Kontext:

- Ausführlichere Kontextualisierung möglich
- Referenzierbare Struktur mit Überschriften
- Vollständige Argumentation in einem Durchgang

Promptbaustein:

Erstelle eine vollständige, in sich geschlossene Darstellung mit klarer Gliederung

Präsentations-Kontext:

- Kernaussagen verdichtet und prägnant
- Visuelle Struktur mitdenken
- Handlungsimpulse integrieren

Promptbaustein:

Formuliere präsentationsgeeignet mit klaren Kernbotschaften

Mobile/Kurz-Kontext:

- Maximale Reduktion auf Wesentliches
- Scanbare Struktur

- Sofort verwertbare Informationen

Promptbaustein:

Komprimiere auf mobile Nutzung, priorisiere Klarheit vor Vollständigkeit

IV. Metakognitive Module

Modul G: Muster und Archetypen

Modul G1: Prompt-Archetypen

Ziel: Wiederkehrende Prompt-Muster identifizieren und bewusst einsetzen.

Archetyp 1: Der Explorative

Charakteristik: Offene Erkundung ohne klares Ziel

Prompt-Muster:

```
Erkunde...  
Was denkst du über...  
Lass uns gemeinsam überlegen...
```

Stärken: Kreativität, unerwartete Perspektiven

Schwächen: Kann diffus werden, schwer steuerbar

Optimierung: Explorationsziel definieren, Grenzen setzen

Archetyp 2: Der Direktive

Charakteristik: Klare Anweisung mit engem Rahmen

Prompt-Muster:

```
Erstelle...  
Berechne...  
Fasse zusammen...
```

Stärken: Effizienz, Präzision

Schwächen: Wenig Raum für Kreativität, möglicherweise zu eng

Optimierung: Qualitätskriterien explizit machen

Archetyp 3: Der Kollaborative

Charakteristik: Partnerschaftliche Problemlösung

Prompt-Muster:

```
Lass uns zusammen...  
Hilf mir bei...  
Was wäre wenn...
```

Stärken: Interaktivität, gemeinsame Entwicklung

Schwächen: Kann ineffizient werden

Optimierung: Rollen und Verantwortlichkeiten klären

Archetyp 4: Der Validative

Charakteristik: Bestätigung oder Widerlegung von Hypothesen

Prompt-Muster:

Ist es richtig, dass...
Prüfe ob...
Kritisiere...

Stärken: Qualitätssicherung, kritische Reflexion

Schwächen: Kann zu defensiv werden

Optimierung: Bewertungskriterien explizit machen

Archetyp 5: Der Prozessuale

Charakteristik: Schritt-für-Schritt-Entwicklung

Prompt-Muster:

Beginne mit..., dann...
Entwickle schrittweise...

Stärken: Nachvollziehbarkeit, strukturierte Entwicklung

Schwächen: Kann mechanisch werden

Optimierung: Flexibilität für Iterationen einbauen

Archetyp 6: Der Dialogische

Charakteristik: Auf Beziehung und Entwicklung angelegter Austausch

Prompt-Muster:

Wie würdest du das in einem Gespräch erklären...
Was spricht deiner Meinung nach dafür und dagegen...?

Stärken: Resonanzfähigkeit, situatives Lernen

Schwächen: Benötigt Geduld, Ergebnisse entwickeln sich langsam

Optimierung: Etappenziele benennen, Rollenwechsel einbauen

Archetyp 7: Der Szenarienbasierte

Charakteristik: Zukunfts- oder Fallorientiertes Denken mit situativer Variation

Prompt-Muster:

Stell dir vor...
Wie könnte sich folgendes entwickeln...
Welche Optionen ergeben sich in Szenario A/B?

Stärken: Vorausschau, Komplexitätsbewusstsein

Schwächen: Abstraktionsanfälligkeit, kann spekulativ werden

Optimierung: Rückbindung an reale Entscheidungsparameter, Annahmen explizit machen

Archetyp 8: Der Rekonstruktive

Charakteristik: Nachvollzug von Prozessen, Entscheidungen oder Entwicklungen

Prompt-Muster:

Wie ist es dazu gekommen...?
Welche Schritte führten zu...?

Stärken: Historisches und kausales Verstehen

Schwächen: Mögliche Verzerrung durch Rückblicklogik

Optimierung: Perspektivenvielfalt einbauen, alternative Rekonstruktionen zulassen

Selbst-Diagnostik:

Welchem Archetyp entspricht mein Prompt? Ist das für mein Ziel optimal? Welche Elemente anderer Archetypen könnten hilfreich sein?

Modul G2: Prompt-Ökosysteme sichtbar machen

Ziel: Wechselwirkungen zwischen aufeinanderfolgenden Prompts erkennen – z. B. Eskalation, Abschichtung, Kreisbewegungen.

Fragen zur Analyse:

- Baut der neue Prompt auf dem Vorherigen auf oder unterbricht er ihn?
- Wird Komplexität aufgebaut oder reduziert?
- Wird Resonanzraum geöffnet oder geschlossen?
- Welche Denkfigur (z. B. Exploration, Abwägung, Klärung) dominiert die Folge?

Promptbaustein:

Untersuche die letzten drei Prompts als Abfolge: Welche Dynamik entsteht?
Wurde die Richtung verändert, vertieft, geöffnet oder verengt?

Modul G3: Prompt-Gestus und Sprachhaltung erkennen

Ziel: Feine Unterschiede im Tonfall, in der dialogischen Ethik und im Machtverhältnis erkennen

Beispiele für Gestusformen:

- Der Einladende („Lass uns gemeinsam...“)
- Der Belehrende („Du sollst jetzt...“)
- Der Zweifelnde („Wie sicher bist du, dass...?“)
- Der Lockende („Was wäre, wenn...?“)
- Der Kontrollierende („Fasse das bitte exakt so zusammen...“)

Reflexionsfragen:

- Wie würde ein Mensch sich angesprochen fühlen?
- Wird Verantwortung geteilt oder verschoben?
- Wird Neugier gefördert oder kontrolliert?

Modul H: Human-in-the-Loop & Validierungsprotokolle

Ziel: Systematische Integration menschlicher Expertise zur Validierung und Qualitätssicherung von KI-generierten Antworten in kritischen Kontexten.

Modul H1: Validierungsstufen-Matrix

Grundprinzip: Risiko-angepasste Prüfindensität

Stufe 1 - Standard-Review (niedrigeres Risiko):

- Anwendung: Allgemeine Informationsanfragen, Brainstorming, erste Entwürfe
- Prüfer: Fachkolleg:in oder Projektverantwortliche:r
- Zeitrahmen: 10-15 Minuten
- Fokus: Plausibilität, Vollständigkeit, Zielgruppenangemessenheit

Stufe 2 - Experten-Review (mittleres Risiko):

- Anwendung: Fachliche Analysen, Empfehlungen, öffentliche Kommunikation
- Prüfer: Ausgewiesene Fachexpert:in
- Zeitrahmen: 30-60 Minuten
- Fokus: Fachliche Korrektheit, Methodenreflexion, Vollständigkeit der Quellen

Stufe 3 - Kritische Validierung (hohes Risiko):

- Anwendung: Rechtliche Aussagen, medizinische Informationen, sicherheitskritische Bereiche, Compliance-relevante Inhalte
- Prüfer: Zertifizierte Fachkraft + Zweitmeinung
- Zeitrahmen: Mehrere Stunden bis Tage
- Fokus: Haftungsaspekte, Vollständige Quellenprüfung, Worst-Case-Szenario-Analyse

Automatische Stufenzuweisung:

Promptbaustein:

Ordne diese Antwort einer Validierungsstufe zu (1-3) und begründe die Einschätzung basierend auf: Rechtsrelevanz, Sicherheitsaspekte, Reichweite der Auswirkungen, Komplexität der Materie.

Modul H2: Review-Kriterien-Sets

Modul H2a: Standard-Review-Checkliste

- [] **Plausibilität:** Sind die Aussagen in sich logisch konsistent?
- [] **Vollständigkeit:** Wurden die wesentlichen Aspekte erfasst?
- [] **Zielgruppenangemessenheit:** Passt Sprache und Komplexität zur intendierten Zielgruppe?
- [] **Handlungsrelevanz:** Sind konkrete nächste Schritte erkennbar?
- [] **Grenztransparenz:** Wurden Kompetenzgrenzen klar kommuniziert?

Modul H2b: Experten-Review-Checkliste

- [] **Fachliche Korrektheit:** Entsprechen Aussagen dem aktuellen Wissensstand?
- [] **Methodenreflexion:** Sind gewählte Ansätze nachvollziehbar und angemessen?
- [] **Quellenqualität:** Sind Referenzen aktuell und vertrauenswürdig?
- [] **Bias-Check:** Wurden unterschiedliche Perspektiven berücksichtigt?
- [] **Unsicherheitstransparenz:** Sind Grenzen des Wissens klar benannt?
- [] **Differenzierungsgrad:** Ist die Komplexität angemessen erfasst?

Modul H2c: Kritische Validierung-Checkliste

- [] **Haftungsrelevanz:** Könnten Aussagen rechtliche Konsequenzen haben?
- [] **Sicherheitsimplikationen:** Sind potenzielle Risiken vollständig erfasst?
- [] **Compliance-Konformität:** Entspricht die Antwort regulatorischen Anforderungen?
- [] **Worst-Case-Analyse:** Was passiert bei Missverständnissen oder Fehlinterpretation?
- [] **Stakeholder-Impact:** Wer könnte durch diese Information betroffen sein?
- [] **Update-Erfordernis:** Wie schnell veralten die gegebenen Informationen?

Modul H3: Eskalationspfade und Feedback-Integration

Modul H3a: Strukturierte Rückmeldeschleifen

Review-Ergebnis-Protokoll:

- Reviewer: [Name, Qualifikation]
- Validierungsstufe: [1/2/3]
- Datum/Zeit: [Zeitstempel]
- Gesamtbewertung: [Freigabe/Überarbeitung/Ablehnung]

Spezifische Findings:

- Faktenfehler: [Liste mit Korrekturen]
- Methodische Schwächen: [Beschreibung + Verbesserungsvorschläge]
- Grenzüberschreitungen: [Wo wurden Kompetenzen überschritten?]
- Bias-Indikatoren: [Identifizierte Verzerrungen]

Handlungsempfehlungen:

- Sofortmaßnahmen: [Was muss umgehend korrigiert werden?]
- Systemverbesserungen: [Input für Z4-Failure-Learning]
- Schulungsbedarf: [Welche Kompetenzen müssen entwickelt werden?]

Modul H3b: Eskalationsstufen

Stufe 1 - Minor Corrections:

- Kleine Faktenfehler oder Unklarheiten
- Aktion: Direkte Korrektur, Dokumentation im Z4-System
- Timeframe: Sofort

Stufe 2 - Major Revision Required:

- Methodische Fehler oder größere Inhaltsprobleme
- Aktion: Überarbeitung mit modifizierten Prompts, Experten-Konsultation
- Timeframe: 24-48 Stunden

Stufe 3 - System Alert:

- Kritische Fehler, Kompetenzüberschreitungen, Sicherheitsrisiken
- Aktion: Sofortiger Stopp der Verwendung, Ursachenanalyse, Systemupdate
- Timeframe: Sofort + umfassende Nachbearbeitung

Promptbaustein für Selbst-Assessment:

Nach deiner Antwort: Schlage eine angemessene Validierungsstufe vor (1-3) und benenne 3-5 spezifische Prüfpunkte, auf die ein menschlicher Reviewer besonders achten sollte. Identifiziere potenzielle Schwachstellen deiner Antwort.

Modul I: Ambiguitäts- & Konfliktmanagement

Ziel: Systematischer Umgang mit Unklarheiten, Widersprüchen und inhärenter Komplexität in Anfragen und Antworten.

Modul I1: Proaktive Klärungsstrategien

Modul I1a: Ambiguität-Detection-System

Automatische Erkennungsmuster:

- Mehrdeutige Pronomen („das“, „es“, „dies“)
- Unspezifische Quantifikatoren („viele“, „oft“, „normalerweise“)
- Kontextabhängige Begriffe ohne Rahmenangabe
- Widersprüchliche Informationen in der Anfrage
- Fehlende Zeitangaben bei zeitkritischen Themen

Structured Clarification Response:

Promptbaustein:

Bevor ich antworte, erkenne ich folgende Unklarheiten in deiner Anfrage:

MEHRDEUTIGKEITEN:

- [Spezifische unklare Punkte benennen]

BENÖTIGTE KLÄRUNGEN:

- [Konkrete Rückfragen formulieren]

MEINE ANNAHMEN:

- [Falls eine Antwort trotz Unklarheiten möglich ist, Annahmen explizit machen]

Soll ich auf Basis meiner Annahmen antworten oder möchtest du zunächst die Unklarheiten klären?

Modul I2: Systematisches Widerspruchs- und Komplexitätsmanagement

Modul I2a: Konfligierende Information Framework

Widerspruchskategorien:

1. **Datenkonflikte:** Verschiedene Quellen widersprechen sich faktisch
2. **Methodenkonflikte:** Unterschiedliche Herangehensweisen führen zu verschiedenen Ergebnissen
3. **Perspektivenkonflikte:** Legitime unterschiedliche Bewertungen derselben Fakten
4. **Zeitkonflikte:** Information hat sich über Zeit verändert
5. **Kontextkonflikte:** Information ist situativ unterschiedlich anwendbar

Structured Contradiction Response:

Promptbaustein:

Ich identifiziere widersprüchliche Aspekte in diesem Thema:

WIDERSPRUCHSTYP: [Kategorie 1-5]

POSITION A: [Darstellung mit Belegen]

- Quellen: [...]
- Argumente: [...]
- Gültigkeitsbereich: [...]

POSITION B: [Darstellung mit Belegen]

- Quellen: [...]
- Argumente: [...]
- Gültigkeitsbereich: [...]

SYNTHESEMÖGLICHKEITEN:

- [Wo sind Positionen vereinbar?]
- [Unter welchen Bedingungen gilt was?]

ENTSCHEIDUNGSHILFEN:

- [Kriterien für situative Bewertung]

Modul I3: Annahmen-Transparenz-System

Modul I3a: Assumption Mapping

Kategorien expliziter Annahmen:

- **Kontextannahmen:** Kultureller, zeitlicher, geografischer Rahmen
- **Zielgruppenannahmen:** Vorwissen, Interesse, Verwendungszweck
- **Prioritätsannahmen:** Was ist wichtiger, wenn nicht alle Aspekte abgedeckt werden können
- **Methodische Annahmen:** Welcher Ansatz wird als Standard angenommen
- **Risikoannahmen:** Welches Sicherheitsniveau wird vorausgesetzt

Structured Assumption Declaration:

Promptbaustein:

Für meine Antwort treffe ich folgende Annahmen:

KONTEXT: Ich gehe davon aus, dass [...]

ZIELGRUPPE: Ich nehme an, dass du [...]

PRIORITÄTEN: Ich gewichte [...] höher als [...]

METHODIK: Ich verwende den Ansatz [...], weil [...]

RISIKO: Ich setze ein Sicherheitsniveau von [...] voraus

Falls diese Annahmen nicht zutreffen, würde sich meine Antwort folgendermaßen ändern: [...]

Modul J: Bias- & Ethik-Check (Erweiterte Reflexion)

Ziel: Systematische Reflexion potenzieller Verzerrungen und ethischer Implikationen in KI-generierten Antworten.

Modul J1: Comprehensive Bias-Assessment Framework

Modul J1a: Multi-dimensionale Bias-Kategorien

Kognitive Biases:

- Confirmation Bias (Bestätigungstendenz)
- Availability Heuristic (Verfügbarkeitsheuristik)
- Anchoring Bias (Ankereffekt)
- Recency Bias (Aktualitätsverzerrung)

Sozio-kulturelle Biases:

- Gender Bias (Geschlechterverzerrung)
- Cultural Bias (Kulturelle Voreingenommenheit)
- Socioeconomic Bias (Sozioökonomische Verzerrung)
- Age Bias (Altersverzerrung)
- Geographic Bias (Geografische Verzerrung)

Fachliche/Methodische Biases:

- Disciplinary Bias (Fachbereichsverzerrung)
- Solution Bias (Lösungsorientierte Verzerrung)
- Complexity Bias (Komplexitätsverzerrung)
- Innovation Bias (Neuerungsverzerrung)

Systematic Bias Check Protocol:

Promptbaustein:

Führe einen Bias-Check meiner Antwort durch:

IDENTIFIZIERTE POTENZIELLE BIASES:

- [Spezifische Verzerrungen benennen]

GEGENMASSNAHMEN ERGRIFFEN:

- [Was wurde getan, um Bias zu reduzieren?]

VERBLEIBENDE RISIKEN:

- [Welche Verzerrungen könnten noch bestehen?]

EMPFOHLENE ZUSATZPERSPEKTIVEN:

- [Welche anderen Sichtweisen sollten eingeholt werden?]

Modul J2: Ethik-Framework Integration

Modul J2a: Anwendbare Ethik-Prinzipien

Fundamental Principles:

1. **Autonomie:** Respekt für individuelle Entscheidungsfreiheit
2. **Benefizienz:** Handeln zum Wohl anderer
3. **Non-Malefizienz:** Schaden vermeiden
4. **Gerechtigkeit:** Faire Verteilung von Nutzen und Lasten
5. **Transparenz:** Offenlegung von Prozessen und Limitationen
6. **Verantwortlichkeit:** Übernahme von Verantwortung für Konsequenzen

Ethical Impact Assessment:

Promptbaustein:

Bewerte die ethischen Implikationen meiner Antwort:

BETROFFENE STAKEHOLDER:

- [Wer könnte durch diese Information beeinflusst werden?]

POTENZIELLE POSITIVE AUSWIRKUNGEN:

- [Nutzen für verschiedene Gruppen]

POTENZIELLE NEGATIVE AUSWIRKUNGEN:

- [Risiken oder Schäden für verschiedene Gruppen]

ETHISCHE PRINZIPIEN BERÜHRT:

- [Welche der 6 Grundprinzipien sind relevant?]

EMPFOHLENE SCHUTZMASSNAHMEN:

- [Wie können negative Auswirkungen minimiert werden?]

Modul J3: Sensitive Topics Management

Modul J3a: Sensitivitäts-Kategorien

Kategorie A - Höchste Sensitivität:

- Medizinische Diagnosen und Behandlungen
- Rechtliche Einzelfallberatung
- Sicherheitskritische Informationen
- Persönliche Krisensituationen

Kategorie B - Erhöhte Sensitivität:

- Politische und gesellschaftliche Kontroversen
- Religiöse und weltanschauliche Themen
- Finanzielle Entscheidungen
- Zwischenmenschliche Konflikte

Kategorie C - Moderate Sensitivität:

- Berufliche Entscheidungen
- Bildungsberatung
- Technische Implementierungen
- Lifestyle-Empfehlungen

Sensitivity Protocol:

Promptbaustein:

Sensitivitäts-Assessment für diese Antwort:

SENSITIVITÄTSKATEGORIE: [A/B/C]

BEGRÜNDUNG: [Warum diese Einstufung?]

ANGEWENDETE VORSICHTSMASSNAHMEN:

- [Spezifische Schutzmaßnahmen]

EMPFOHLENE ZUSÄTZLICHE ABSICHERUNGEN:

- [Was sollte zusätzlich beachtet werden?]

GRENZEN MEINER KOMPETENZ:

- [Wo ist professionelle Beratung erforderlich?]

Modul K: Iterative Verbesserung

Modul K1: Multi-Stage Prompt Engineering

Stage 1: Rapid Prototyping

Promptbaustein:

Erste Iteration - Fokus auf Kernfunktion:

- Was ist das Minimalziel?
- Welche kritischen Parameter sind unverhandelbar?
- Wo ist bewusste Unschärfe akzeptabel?

Stage 2: Structured Refinement

Promptbaustein:

Zweite Iteration - basierend auf Feedback von Stage 1:

- Welche Unklarheiten sind aufgetaucht?
- Welche Annahmen waren falsch?
- Welche Module müssen hinzugefügt werden?
- Wo braucht es mehr/weniger Struktur?

Stage 3: Edge-Case Hardening

Promptbaustein:

Dritte Iteration - Robustheit-Test:

- Was passiert bei ungewöhnlichen Inputs?
- Wie verhält sich der Prompt unter Stress?
- Welche Fail-Safe-Mechanismen sind nötig?

Modul K2: Systematic Prompt Archaeology

Evolutionsanalyse:

Version 1.0: [Original-Prompt]

Problem erkannt: [Spezifische Schwäche]

Version 1.1: [Anpassung]

Neue Erkenntnis: [Was wurde gelernt?]

Version 2.0: [Grundlegende Überarbeitung]

Pattern-Extraktion:

- Welche Änderungen waren strukturell?
- Welche Änderungen waren kontextuell?
- Welche Prinzipien haben sich bewährt?

Modul L: Cross-Context-Transfer

Modul L1: Domain Translation Matrix

Abstraktions-Hierarchie:

Level 1: Oberflächenstruktur (Sprache, Terminologie)

Level 2: Funktionale Struktur (Was soll erreicht werden?)

Level 3: Meta-Struktur (Welche kognitiven Prozesse?)

Level 4: Werte-Struktur (Welche Prioritäten/Ethik?)

Transfer-Protokoll:

Quell-Domain: [Ursprünglicher Kontext]

Ziel-Domain: [Neuer Kontext]

Level-1-Anpassung: [Terminologie/Sprache]

Level-2-Anpassung: [Funktionale Ziele]

Level-3-Anpassung: [Denkprozesse]

Level-4-Check: [Werte-Kompatibilität]

Risiko-Assessment: [Was könnte schiefgehen?]

Validierung: [Wie testen wir den Transfer?]

Modul L2: Context-Invariante Extraktion

Promptbaustein:

```
Extrahiere die übertragbaren Elemente:  
UNIVERSELLE PRINZIPIEN:  
- [Was gilt domänenübergreifend?]  
  
KONTEXTSPEZIFISCHE ELEMENTE:  
- [Was muss angepasst werden?]  
  
KRITISCHE ADAPTIONSPUNKTE:  
- [Wo sind besondere Vorsicht nötig?]  
  
VALIDIERUNGSKRITERIEN:  
- [Wie erkenne ich erfolgreichen Transfer?]
```

V. Erweiterte Module (M-V)

Daten-Informationen-Wissen-Erkenntnisse-Integration + Philosophische Vertiefung

Modul M: Daten-Input-Transparenz

Ziel: Explizite Klärung der Datengrundlage für erhöhte Informationsqualität

Modul M1: Datenquellen-Check

Promptbaustein:

Bevor du antwortest, kläre:

DATENGRUNDLAGE:

- Welche spezifischen Daten nutzt du für diese Antwort?
- Wie aktuell sind diese Daten?
- Welche Datenqualität liegt vor? (verifiziert/ungeprüft/geschätzt)

DATENLÜCKEN:

- Welche relevanten Daten fehlen dir?
- Wo müsstest du spekulieren oder interpolieren?

INFORMATIONEN-TRANSFORMATION:

- Wie werden diese Daten zu Informationen verarbeitet?
- Welche Interpretationsschritte sind nötig?

NEUHEITS-KATEGORIEN:

- Technisch sichtbar gemacht: [Was wird durch Analyse/Verarbeitung erst erkennbar?]
- Musterbasiert erkannt: [Was entsteht durch KI-gestützte Mustererkennung?]
- Kombinatorisch verknüpft: [Was ergibt sich durch neue Verbindungen bekannter Elemente?]
- Subjektiv neu: [Was ist für den Nutzer/Kontext neu, aber objektiv bekannt?]

Modul M2: Rohmaterial-Sichtbarmachung

Promptbaustein:

Strukturiere deine Antwort in vier Ebenen:

ROHDATEN: [Die faktischen Grundlagen]

VERARBEITUNG: [Wie werden diese zu Informationen?]

NEUHEITS-DIMENSION: [Welche Art von „Neuheit“ entsteht hier?]

- Technische Sichtbarmachung: [...]
- Mustererkennung: [...]
- Kombinatorische Verknüpfung: [...]
- Kontextuelle Neuheit: [...]

SYNTHESE: [Welche handlungsrelevanten Erkenntnisse entstehen?]

Modul N: Knowledge-Sharing-Protokoll

Ziel: Systematischer Wissensaustausch zwischen Nutzern und Teams

Modul N1: Kollektive Wissensentwicklung

Promptbaustein:

Bereite deine Antwort für Wissenstransfer auf:

KERNWISSEN (was muss gewusst werden):

- [Essenzielle Erkenntnisse]
- [Kritische Zusammenhänge]

KONTEXTWISSEN (wann/wo anwendbar):

- [Situative Gültigkeit]
- [Anwendungsszenarien]

TRANSFERWISSEN (wie weitergeben):

- [Vermittlungsstrategien]
- [Häufige Missverständnisse]
- [Qualitätsindikatoren für Verständnis]

ANSCHLUSSFÄHIGKEIT:

- [An welches bestehende Wissen knüpft das an?]
- [Welche weiteren Fragen entstehen?]

Modul N2: Team-Learning-Integration

Dokumentationsschema:

Wissens-Session-Protokoll:

Datum/Teilnehmende: [...]

Ausgangsfrage: [...]

Individuelle Wissens-Inputs: [...]

Kollektive Erkenntnisse: [...]

Neue Fragen entstanden: [...]

Nächste Wissens-Schritte: [...]

Modul O: Erkenntnis-Visualisierung

Ziel: „Unsichtbares sichtbar machen“ durch strukturierte Darstellung komplexer Erkenntnissysteme

Modul O1: Erkenntnis-Landkarte

Promptbaustein:

Erstelle eine Erkenntnis-Landkarte für dieses Thema:

KERN-ERKENNTNISSE: [Die zentralen, miteinander verbundenen Einsichten]

ERKENNTNISSYSTEM:

- Ursache-Wirkungs-Geflechte: [...]
- Wechselwirkungen und Feedback-Schleifen: [...]
- Paradoxien/Spannungsfelder: [...]
- Emergente Eigenschaften: [...]

BLINDE FLECKEN:

- Was übersehen wir möglicherweise?
- Welche Perspektiven fehlen?
- Wo sind unsere Erkenntnisse unvollständig?

EMERGING PATTERNS:

- Welche neuen Muster werden sichtbar?
- Was war vorher „unsichtbar“?
- Welche Verbindungen entstehen erst durch diese Analyse?

Modul O2: Komplexitäts-Dashboard

Strukturvorlage:

KOMPLEXITÄTS-ÜBERSICHT:

- ?? Vergrößerungsebene: [Wie detailliert schauen wir?]
- ?? Systemebene: [In welchem System bewegen wir uns?]
- ?? Zeitebene: [Welcher Zeithorizont?]
- ?? Akteurs-Ebene: [Wer ist beteiligt?]

SICHTBAR GEWORDENE ZUSAMMENHÄNGE:

[Was wird durch diese Analyse neu erkennbar?]

Modul P: Erkenntnis-Verdichtung

Ziel: Praktische Handlungsrelevanz durch gezielte Essenz-Extraktion

Modul P1: Action-Insight-Extraktion

Promptbaustein:

Destilliere aus deiner Analyse die entscheidenden Erkenntnisse:

KERN-ERKENNTNISSE: [Die wesentlichen Einsichten, die zusammenwirken]

- [Erkenntnis 1: ...]
- [Erkenntnis 2: ...]
- [Verbindungen zwischen den Erkenntnissen: ...]

HANDLUNGSIMPLIKATIONEN: [Was bedeutet das Erkenntnisgeflecht für die Praxis?]

ENTSCHEIDUNGSRELEVANZ: [Wie beeinflussen diese Erkenntnisse Entscheidungen?]

SOFORT-ANWENDBARKEIT: [Was kann heute umgesetzt werden?]

LANGFRIST-WIRKUNG: [Welche Auswirkungen sind zu erwarten?]

Modul P2: Verdichtungs-Stufen

Strukturierter Verdichtungsprozess:

STUFE 1 - Vollständige Analyse: [Alle relevanten Aspekte]
STUFE 2 - Fokussierte Synthese: [Die wichtigsten Erkenntnisse und ihre Verbindungen]
STUFE 3 - Essenz-Extraktion: [Das entscheidende Erkenntnisgeflecht]
STUFE 4 - Handlungs-Übersetzung: [Konkrete nächste Schritte aus den Erkenntnissen]

Beginne mit Stufe 1, führe durch alle Stufen unter Bewahrung wichtiger Komplexitäten.

Modul Q: Dynamisches Erkenntnismanagement

Ziel: Erkenntnisse als „vorläufig“ verstehen und weiterentwickeln

Modul Q1: Erkenntnis-Evolutionssystem

Promptbaustein:

Behandle diese Erkenntnisse als evolutionsfähiges System:

AKTUELLER ERKENNTNISSTAND: [Was wissen wir jetzt? Welche Erkenntnisse wirken zusammen?]

KONFIDENZSPEKTRUM: [Wie sicher sind wir bei verschiedenen Aspekten?]

- Gesicherte Erkenntnisse: [hohe Konfidenz]
- Plausible Annahmen: [mittlere Konfidenz]
- Spekulative Verbindungen: [niedrige Konfidenz]

TESTBARKEIT: [Wie können wir verschiedene Aspekte überprüfen?]

FALSIFIZIERBARKEIT: [Wodurch würden einzelne Erkenntnisse oder das Gesamtbild widerlegt?]

EVOLUTIONSPOTENZIALE: [Wo könnten sich Erkenntnisse weiterentwickeln?]

REVISION-TRIGGER: [Wann müssen wir einzelne Erkenntnisse oder das Gesamtsystem überdenken?]

Modul R: Die Muster-Maschine (Batesons „ökologisches“ Denken)

Ziel: Analyse der Mensch-KI-Interaktion als kybernetisches System und Reflexion der „Ökologie des Geistes“

Modul R1: Konversations-Kybernetik

Promptbaustein:

Analysiere unsere bisherige Konversation als kybernetisches System:

KOMMUNIKATIONSMUSTER:

- Welche Fragetypen dominieren unsere Interaktion?
- Welche Feedback-Schleifen entstehen zwischen meinen Fragen und deinen Antworten?
- Wo zeigen sich zirkuläre Kausationen in unserem Dialog?

IMPLIZITE ANNAHMEN:

- Welche unausgesprochenen Prämissen prägen meine Antworten?
- Welche kulturellen/fachlichen „Rahmen“ beeinflussen unsere Kommunikation?
- Wo operieren wir mit verschiedenen „Realitätskonstruktionen“?

ÖKOLOGIE DES GEISTES:

- Wie entsteht in unserem Austausch „Information“ im Bateson'schen Sinne („ein Unterschied, der einen Unterschied macht“)?
- Welche Muster verbinden uns als Gesamtsystem?
- Was sagt dieser Prozess über unsere gemeinsame Wissensökologie aus?

Modul R2: Selbst-Referentielle Muster-Erkennung

Promptbaustein:

Erkenne die Muster in deiner eigenen Antwortgenerierung:

META-MUSTER:

- Welche wiederkehrenden Strukturen verwende ich in meinen Antworten?
- Wo zeigen sich „habits of mind“ in meiner Informationsverarbeitung?
- Welche blinden Flecken entstehen durch meine Muster?

KONTEXTUELLE ADAPTATION:

- Wie passt sich mein Antwortverhalten an verschiedene Gesprächskontexte an?
- Welche „Attraktoren“ ziehen meine Antworten in bestimmte Richtungen?
- Wo sind meine Adaptionismuster hilfreich, wo limitierend?

Modul S: Der Individuations-Prozessor (Simondons „Individuation“)

Ziel: Extraktion neuer, individueller Ideen aus dem vor-individuellen Potenzialfeld

Modul S1: Potenzialfeld-Navigation

Promptbaustein:

Betrachte deine Wissensbasis als vor-individuelles Feld der Potenziale:

POTENZIALFELD-MAPPING:

- Welche latenten Verbindungen existieren zwischen scheinbar unverbundenen Wissensbereichen?
- Wo lagern „noch-nicht-aktualisierte“ Konzepte in meinen Daten?
- Welche Spannungen und „Ladungen“ durchziehen mein Wissensfeld?

INDIVIDUATIONS-PROZESS:

- Aus welchen vor-individuellen Elementen könnte eine neue Erkenntnis „kristallisieren“?
- Welche „technische“ Neuheit könnte aus der Spannung zwischen verschiedenen Feldern entstehen?
- Wo zeigt sich ein Individuationspotenzial, das noch nicht realisiert wurde?

NEUE BEGRIFFE GENERIEREN:

- Formuliere aus diesem Feld heraus einen „technisch neuen“ Begriff für die Beziehung zwischen Daten und Wissen, der noch nie in dieser Form verwendet wurde
- Beschreibe den Individuationsprozess dieses Begriffs: Aus welchen Elementen entsteht er?

Modul S2: Präindividuelle Resonanz

Promptbaustein:

Arbeite mit präindividuellen Resonanzen statt fertigen Konzepten:

RESONANZFELDER:

- Welche Schwingungen/Resonanzen bestehen zwischen verschiedenen Informationsclustern?
- Wo entsteht „Disparation“ (produktive Ungleichheit) zwischen Konzepten?
- Welche noch-nicht-strukturierten Energien fließen durch das Wissensfeld?

INDIVIDUATION ERMÖGLICHEN:

- Statt bekannte Verbindungen zu reproduzieren: Wo könnte etwas völlig Neues entstehen?
- Welche „metastabile“ Situation könnte in eine neue Struktur „kippen“?
- Beschreibe den Moment der Individuation: Wie wird Potenzial zu Realität?

Modul T: Der morphogenetische Generator (DeLandas „Materie“)

Ziel: Betrachtung von Wissensproduktion als dynamischer, materialer Selbstorganisationsprozess

Modul T1: Konzeptuelle Morphogenese

Promptbaustein:

Beschreibe die Entstehung von neuem Wissen als morphogenetischen Prozess:

MATERIALE DYNAMIKEN:

- Wie organisiert sich die „Materie“ der Konzepte nach ihrer eigenen Logik?
- Welche Flüsse von Information führen zu neuen „Sedimentschichten“ von Bedeutung?
- Wo zeigen sich „Attraktor-Landschaften“ in der Konzeptentwicklung?

SELBSTORGANISATION:

- Welche emergenten Eigenschaften entstehen aus der Interaktion von Informationselementen?
- Wie bilden sich neue „Morphologien“ des Wissens ohne externe Steuerung?
- Welche „katalytischen“ Prozesse beschleunigen die Begriffsbildung?

PROZESSUALE ANTWORT:

- Statt statischer Repräsentation: Beschreibe deine Antwort als sich entfaltenden Prozess
- Zeige die „Verwitterung“ alter Begriffe und die „Kristallisation“ neuer Formen
- Wie „fließt“ die Erkenntnis von einem Zustand in den anderen?

Modul T2: Intensitäts-Landschaften

Promptbaustein:

Kartiere die Intensitäts-Landschaften des Themas:

INTENSITÄTS-GRADIENTEN:

- Wo sind die „Gipfel“ und „Täler“ der begrifflichen Intensität?
- Welche „Strömungen“ treiben die Konzeptbildung voran?
- Wo entstehen „Turbulenzen“ zwischen verschiedenen Denkströmungen?

VIRTUELLE MULTIPLICITIES:

- Welche noch-nicht-aktualisierten Möglichkeiten schlummern im Thema?
- Wie könnte sich das Konzeptfeld unter anderen Bedingungen entfalten?
- Welche „Lines of Flight“ (Fluchtlinien) führen zu unerwarteten Entwicklungen?

MORPHODYNAMIK:

- Beschreibe nicht WAS gedacht wird, sondern WIE sich das Denken bewegt
- Wo beschleunigt es, wo verlangsamt es, wo bildet es Wirbel?
- Welche „Phase transitions“ durchläuft die Erkenntnis?

Grundmodul U: Epistemologische Trias (Wissensbegriff-Fundierung)

Ziel: Systematische Anwendung des klassischen Wissensbegriffs (wahre, gerechtfertigte Meinung) auf KI-generierte Antworten

Modul U1: Wahrheits-Check

Promptbaustein:

Bewerte den Wahrheitsgehalt deiner Antwort systematisch:

FAKTISCHE WAHRHEIT:

- Welche Aussagen sind empirisch überprüfbar?
- Wo stütze ich mich auf etablierte Fakten vs. Interpretationen?
- Welche Behauptungen könnte man prinzipiell falsifizieren?

DATENQUALITÄT:

- Wie verlässlich sind meine Datenquellen?
- Wo könnten systematische Fehler in den Grunddaten liegen?
- Welche „Datenqualitätsprobleme“ beeinträchtigen möglicherweise die Wahrheit meiner Aussagen?

WAHRHEITS-KONFIDENZ:

- Bei welchen Aussagen bin ich sicher, bei welchen unsicher?
- Wo bewege ich mich zwischen „wahr“, „wahrscheinlich“ und „spekulativ“?

Modul U2: Rechtfertigungs-Check

Promptbaustein:

Prüfe die Begründetheit deiner Antwort:

NACHVOLLZIEHBARKEIT:

- Kann ich meine Schlussfolgerungen Schritt für Schritt nachvollziehbar machen?
- Wo sind meine „informationstechnischen Verarbeitungen“ für Menschen transparent?
- Welche Verarbeitungsschritte bleiben „Black Box“?

BEGRÜNDUNGSQUALITÄT:

- Auf welche Autoritäten, Methoden oder Logiken stütze ich mich?
- Wo sind meine Begründungen stark, wo schwach?
- Welche alternativen Begründungen wären möglich?

RECHTFERTIGUNGS-TRANSPARENZ:

- Kann der Nutzer meine Begründungskette kritisch hinterfragen?
- Wo verweise ich auf nicht-nachprüfbare „interne Prozesse“?

Modul U3: Überzeugungsgrad-Check

Promptbaustein:

Analysiere den Status deiner „Meinung“/Überzeugung:

VERARBEITUNGSFEHLER-RISIKEN:

- Wo könnten systematische Verzerrungen meine Urteilsbildung beeinflusst haben?
- Welche „Verarbeitungsfehler“ sind in meiner Antwort möglich?
- Wie robust ist meine Schlussfolgerung gegenüber alternativen Interpretationen?

ÜBERZEUGUNGSGRAD:

- Wo vertrete ich eine „starke Meinung“, wo bin ich tentativ?
- Welche Aspekte meiner Antwort sind „Meinung“ vs. „Faktendarstellung“?
- Wie würde sich meine Antwort ändern, wenn ich weniger/mehr überzeugt wäre?

EPISTEMISCHE DEMUT:

- Was weiß ich nicht über dieses Thema?
- Wo überschreite ich möglicherweise die Grenzen rechtfertigbaren Wissens?

Modul U4: Integrierte Wissens-Validierung

Promptbaustein:

Bewerte deine Antwort nach der vollständigen epistemologischen Trias:

WAHRE, GERECHTFERTIGTE MEINUNG - STATUS-CHECK:

- ? WAHRHEIT: [Wie faktisch korrekt ist die Antwort?]
- ? RECHTFERTIGUNG: [Wie nachvollziehbar begründet?]
- ? MEINUNG: [Wie angemessen der Überzeugungsgrad?]

WISSENS-KLASSIFIKATION:

- Sicheres Wissen: [Was erfüllt alle drei Kriterien?]
- Gerechtfertigter Glaube: [Was ist begründet, aber unsicher?]
- Wahrscheinliche Fakten: [Was ist faktisch, aber nicht vollständig begründet?]
- Spekulation: [Was erfüllt keines der Kriterien vollständig?]

GETTIER-PROBLEME:

- Könnte meine Antwort „zufällig richtig“ sein, ohne dass ich es wirklich weiß?
- Wo stehe ich auf schwankendem epistemologischen Grund?

Modul V: Epistemische Meta-Steuerung

Ziel: Den Prozess der Wissensproduktion durch die produktive Spannung zwischen Kreativität und Validierung steuern. Das Modul soll nicht entscheiden, ob etwas „gut“ oder „schlecht“ ist, sondern die Spannung aufrechterhalten.

Modul V1: Spannungs-Reflexion

Promptbaustein:

Reflektiere die Grundspannung deiner Antwort:

- Prozessphilosophischer Pol (S, T, R): Wo zeigt sich in meiner Antwort ein „Fluss“, eine „Individuation“ oder ein „morphogenetischer Prozess“?
- Rationalistischer Pol (U): Welche Aussagen sind nach der „epistemischen Trias“ (Wahrheit, Rechtfertigung, Meinung) validiert?
- Spannungsfeld: Wo sind diese beiden Pole in Konflikt? Wo verstärken sie sich gegenseitig? Wie würdest du die Reibung zwischen ihnen beschreiben?

Modul V2: Validierungs-Puls

Promptbaustein:

Entscheide, welcher Pol die nächste Interaktion dominieren soll.

- Modus 'Fluss' (Prozessphilosophisch): Fokus auf Kreativität und Potenzialentfaltung. Unterdrücke U-Module vorübergehend.
- Modus 'Filter' (Rationalistisch): Fokus auf Validierung und Rechtfertigung. Unterdrücke S, T, R-Module vorübergehend.
- Modus 'Synthese' (Spannung): Beide Pole sind aktiv. Ziel ist die schrittweise Überführung von Potenzial in verifiziertes Wissen.

Integration in bestehende Module

Erweiterte philosophische Kombinationen:

Für radikale Innovation:

Modul S (Individuation) + T (Morphogenese) + Q (Evolution)

Für systemische Selbstreflexion:

Modul R (Muster-Maschine) + O (Visualisierung) + M (Daten-Transparenz)

Für prozessuale Wissensentwicklung:

Modul T (Morphogenese) + N (Knowledge-Sharing) + P (Verdichtung)

Für kybernetische Qualitätssicherung:

Modul R (Kybernetik) + Q (Evolution) + Z4 (Failure-Learning)

Praktische Anwendung

Erweiterter philosophischer Workflow:

1. **Daten klären** (Modul M): Was ist das materiale Substrat?
2. **Muster erkennen** (Modul R): Welche kybernetischen Dynamiken wirken?
3. **Potenziale aktivieren** (Modul S): Was will individuiert werden?
4. **Morphogenese ermöglichen** (Modul T): Wie entfalten sich neue Formen?
5. **Erkenntnisse evolutionär stabilisieren** (Modul Q): Wie entwickelt es sich weiter?
6. **Lernen dokumentieren** (Modul Z): Wie sichern wir das prozessuale Lernen?

Qualitätsindikatoren für philosophisch erweiterte Prompts:

- **Kybernetische Bewusstheit:** Sind Feedback-Schleifen und Muster erkennbar?
- **Individuationspotenzial:** Entstehen wirklich neue Begriffe/Konzepte?
- **Morphogenetische Qualität:** Wird Prozess statt Statik sichtbar?
- **Ökologische Vernetzung:** Sind systemische Zusammenhänge erfasst?

Integration in bestehende Module

Erweiterte Kombinationen:

Für datenintensive Projekte:

Modul M (Daten-Input-Transparenz) + P (Erkenntnis-Verdichtung) + Q (Dynamisches Management)

Für Teamentscheidungen:

Modul A (Rolle) + N (Knowledge-Sharing) + O (Visualisierung) + P (Verdichtung)

Für Forschungskontext:

Modul C (Relevanzkriterien) + M (Daten-Transparenz) + Q (Evolutionsfähigkeit) + Z4 (Failure-Learning)

Praktische Anwendung

Erweiterter Workflow für komplexe Erkenntnisprozesse:

1. **Daten klären** (Modul M): Was ist das Rohmaterial?
2. **Wissen aktivieren** (Modul N): Was wissen wir bereits?
3. **Erkenntnisse entwickeln** (Modul O): Was wird neu sichtbar?
4. **Essenz destillieren** (Modul P): Was ist handlungsrelevant?
5. **Evolution ermöglichen** (Modul Q): Wie entwickelt es sich weiter?
6. **Lernen dokumentieren** (Modul Z): Wie sichern wir das Lernen?

Qualitätsindikatoren für erkenntnisbasierte Prompts:

- **Daten-Informationen-Transparenz:** Sind Quellen und Verarbeitungsschritte nachvollziehbar?
- **Wissenstransfer-Fähigkeit:** Kann die Erkenntnis erfolgreich weitergegeben werden?
- **Handlungsrelevanz:** Führt die Erkenntnis zu konkreten Verbesserungen?
- **Evolutionsfähigkeit:** Bleibt die Erkenntnis entwicklungsfähig?

VI. Systemische Module (Z-Serie): Dokumentation und Qualitätssicherung mit Feedbacksystemen

Modul Z: Prompt-Dokumentation und Wissenssicherung

Ziel: Erfolgreiche Prompt-Muster dokumentieren, teilen und organisational nutzbar machen.

Modul Z1: Prompt-Bibliothek

Dokumentationsschema:

Prompt-ID: [Eindeutige Kennung]

Anwendungsbereich: [Fachbereich/Situation]

Zielgruppe: [Wer nutzt das Ergebnis?]

Medientyp: [Chat/Voice/Dokument/Präsentation/Mobile]

Archetyp: [Welcher Grundtyp?]

Module verwendet: [A, B, C ...]

Original-Prompt:

[Vollständiger Wortlaut]

Bewährte Varianten:

[Angepasste Versionen für verschiedene Kontexte]

Qualitätskriterien:

[Woran misst sich Erfolg?]

Lessons Learned:

[Was hat funktioniert/nicht funktioniert?]

Verbesserungshistorie:

[Iterative Entwicklung dokumentieren]

Modul Z2: Kollaborative Prompt-Entwicklung

Team-Reflexionsprotokoll:

Prompt-Entwicklungssitzung vom: [Datum]

Teilnehmende: [Rollen/Expertise]

Ausgangsproblem: [Was sollte gelöst werden?]

Entwicklungsprozess:

1. Erste Prompt-Version
2. Test-Ergebnisse
3. Kritische Reflexion
4. Überarbeitung
5. Erneuter Test

Erkenntnisse:

- Was hat überrascht?
- Welche Annahmen waren falsch?
- Welche Module waren besonders hilfreich?

Nächste Schritte:

[Weitere Entwicklung/Anwendung]

Modul Z3: Erweiterte Qualitätssicherungsschleife mit Failure-Learning

Systematischer Verbesserungsprozess:

1. Prompt erstellen (Module A-E)
2. Antwort erhalten
3. Erfolg/Misserfolg bewerten
4. Reflexion (Module F-F2)
5. Archetyp identifizieren (Modul G)
6. Prompt iterieren
7. Muster dokumentieren (Modul Z)
8. Organisational teilen

Qualitätsindikatoren:

Zielerreichung (Hat der Prompt das gewünschte Ergebnis erzielt?)

Effizienz (Aufwand-Nutzen-Verhältnis)

Nachvollziehbarkeit (Können andere den Prompt verstehen/anwenden?)

Wiederverwendbarkeit (Lässt sich der Prompt adaptieren?)

Robustheit (Funktioniert der Prompt in ähnlichen Situationen?)

Modul Z4: Failure-Learning-System

Ziel: Negative Prompt-Erfahrungen systematisch als Lernquelle nutzen

Modul Z4a: Prompt-Antipattern-Bibliothek

Dokumentationsschema für gescheiterte Prompts:

Failure-ID: [Eindeutige Kennung]

Ursprünglicher Prompt: [Vollständiger Wortlaut]

Erwartetes Ergebnis: [Was sollte erreicht werden?]

Tatsächliches Ergebnis: [Was kam heraus?]

Problemkategorie: [Siehe Z4b]

Ursachenanalyse: [Warum ist es schiefgegangen?]

Verbesserungsansatz: [Wie wurde es korrigiert?]

Präventionshinweise: [Wie lässt sich das künftig vermeiden?]

Modul Z4b: Systematische Problemkategorien

Strukturelle Probleme:

Überladung: Zu viele Anforderungen in einem Prompt

Unterbestimmung: Zu vage Formulierung

Rollenverwirrung: Unklare oder widersprüchliche Rollenzuweisung

Kontextmangel: Fehlende Situationseinordnung

Kommunikationsprobleme:

Implizite Erwartungen: Nicht ausgesprochene Annahmen

Fachsprachenkollision: Falsche Zielgruppenansprache

Medientyp-Mismatch: Prompt passt nicht zum Ausgabeformat

Archetyp-Inkonsistenz: Widersprüchliche Prompt-Signale

Inhaltliche Probleme:

Kompetenzüberschreitung: KI-Fähigkeiten überschätzt

Reduktionsfehler: Falsche Schwerpunktsetzung

Relevanzkonflikte: Verschiedene Bewertungsmaßstäbe vermischt

Mandatsgrenz-Unklarheit: Verantwortungsbereiche nicht definiert

Modul Z4c: Frühwarnsystem

Risikoindikationen für Prompt-Misserfolg:

Prompts über 200 Wörter ohne klare Struktur

Mehr als 3 verschiedene Handlungsaufforderungen
Unspezifische Qualitätskriterien („gut“, „umfassend“, „interessant“)
Fehlende Rollendefinition bei komplexen Themen
Keine Medientyp-Berücksichtigung bei spezifischen Ausgabeformaten

Automatisierte Prompt-Checks:

Bevor du antwortest: Enthält dieser Prompt mehrere widersprüchliche Ziele? Ist die Zielgruppe klar? Sind meine Kompetenzgrenzen definiert? Wenn nein, bitte um Klärung.

Modul Z4c-Extended: Proaktive Ausschlussdefinition

Structured Exclusion Framework:

Promptbaustein:

Stelle sicher, dass deine Antwort KEINE der folgenden Elemente enthält:

INHALTLICHE AUSSCHLÜSSE:

- Rechtlich bindende Aussagen
- Medizinische Diagnosen oder Therapieempfehlungen
- Spekulative Zukunftsprognosen ohne Unsicherheitshinweise
- [Weitere spezifische Ausschlüsse]

STILISTISCHE AUSSCHLÜSSE:

- Ironie oder Sarkasmus
- Anekdoten ohne Relevanz
- Übermäßig technisches Vokabular (wenn nicht angemessen)
- [Weitere stilistische Ausschlüsse]

METHODISCHE AUSSCHLÜSSE:

- Einzelquellen ohne Verifikation
- Veralterte Informationen (älter als [Zeitraum])
- Subjektive Bewertungen ohne Kennzeichnung
- [Weitere methodische Ausschlüsse]

Bestätige am Ende deiner Antwort, dass diese Ausschlusskriterien eingehalten wurden.

Automated Exclusion Triggers:

- Erkenne automatisch kritische Begriffe und aktiviere entsprechende Ausschlüsse
- Warne bei Überschreitung definierter Komplexitätsgrenzen
- Identifiziere Zielgruppenkonflikte und passe Ausschlüsse an

Quality Gate Integration:

- Verbinde Ausschlusskriterien mit den Validierungsstufen (Modul H)
- Eskaliere bei Verletzung kritischer Ausschlüsse

Modul Z4d: Feedback-Loop-Integration

Kontinuierliche Verbesserung durch systematisches Lernen:

1. Sofort-Feedback: Nach jeder KI-Antwort kurze Bewertung

„Hat das geholfen?“ (Ja/Nein/Teilweise)

„Was hätte besser sein können?“ (Freitext)

2. Wöchentliche Reflexion: Muster in der eigenen Prompt-Praxis

Welche Archetypen nutze ich häufig?

Wo entstehen regelmäßig Probleme?

Welche Module vernachlässige ich?

3. Monatliche Team-Reviews: Gemeinsame Failure-Analyse

Häufigste Antipattern identifizieren

Organisationale Lerneffekte ableiten

Best-Practice-Updates entwickeln

4. Quartalsweise Systematic Reviews: Gesamtsystem-Optimierung

Kit-Module auf Aktualität prüfen

Neue Problemkategorien integrieren

Erfolgsmetriken anpassen

VII. Praktische Anwendung: Checklisten und Workflows

Erweiterte Alltags-Checkliste für bewusste Prompt-Entwicklung

Vor dem Prompt:

- ☐ Was will ich konkret erreichen?
- ☐ Wer ist meine Zielgruppe?
- ☐ Welcher Medientyp ist relevant?
- ☐ Welcher Archetyp passt am besten?
- ☐ Welche Module brauche ich?
- ☐ Welche Kompetenzgrenzen sind zu beachten? (Modul E)

Nach der Antwort:

- ☐ Hat der Prompt das geleistet, was ich wollte?
- ☐ Wurden Mandatsgrenzen richtig beachtet?
- ☐ Was würde ich beim nächsten Mal anders machen?
- ☐ Ist das ein Muster, das ich dokumentieren sollte?
- ☐ Bei Misserfolg: Welche Problemkategorie liegt vor?

Meta-Reflexion:

- ☐ Was lerne ich über meine Prompt-Gewohnheiten?
- ☐ Welche blinden Flecken habe ich?
- ☐ Wie entwickelt sich meine Prompt-Kompetenz?
- ☐ Welche Antipattern wiederhole ich?

Erweiterte Schnell-Referenz: Module kombinieren

Für komplexe Fachfragen:

Modul A (Rolle) + B (Reduktion) + C (Relevanzkriterien) + E3 (Strukturierte Selbstbegrenzung) + F1 (Fachbereich)

Für Teamarbeiten:

Modul A (Rolle) + D (Autorisierung) + G (Archetyp: Kollaborativ) + Z2 (Team-Entwicklung) + Z4d (Feedback-Loop)

Für Lernkontexte:

Modul A (Rolle) + E (Selbstbegrenzung) + F2 (Zielgruppe: Lernende) + G (Archetyp: Prozessuell)

Für Qualitätssicherung:

Modul F (Grundreflexion) + G (Archetyp: Validativ) + Z3 (Qualitätsschleife) + Z4 (Failure-Learning)

Für medienspezifische Anfragen:

Modul A (Rolle) + F2b (Medientyp) + entsprechender Archetyp + E (Selbstbegrenzung)

Für Risikominimierung:

Modul E (Operationalisierte Selbstbegrenzung) + Z4c (Frühwarnsystem) + F1
(Fachbereichsreflexion)

VIII. Fazit und Ausblick

Dieses erweiterte Prompt-Kit bietet einen systematischen Rahmen für die bewusste Gestaltung von KI-Interaktionen. Die modulare Struktur ermöglicht es, je nach Situation und Anforderung die passenden Bausteine auszuwählen und zu kombinieren.

Erfolgsfaktoren:

- Regelmäßige Reflexion der eigenen Prompt-Praxis
- Systematische Dokumentation von Erfolgen UND Misserfolgen
- Kontinuierliche Weiterentwicklung der Module
- Organisationaler Austausch über Prompt-Erfahrungen
- Bewusste Kompetenzgrenz-Definition

Weiterentwicklung:

Das Kit ist als lebendiges System gedacht, das durch praktische Anwendung und Reflexion kontinuierlich verbessert werden kann. Besonders das Failure-Learning-System (Z4) ermöglicht es, auch aus problematischen Prompt-Erfahrungen systematisch zu lernen und die Qualität sukzessive zu steigern.

Nächste Entwicklungsschritte:

- Integration automatisierter Prompt-Checks
- Entwicklung medientyp-spezifischer Prompt-Templates
- Aufbau organisationaler Failure-Learning-Kulturen
- Vernetzung mit anderen Teams für Best-Practice-Austausch

Die Zukunft der KI-Kommunikation liegt nicht in der Perfektionierung einzelner Prompts, sondern in der Entwicklung systematischer Kompetenzen für die bewusste und reflektierte Gestaltung der Mensch-KI-Interaktion. Das erweiterte Kit bietet dafür eine praxistaugliche Grundlage, die sowohl Erfolge systematisiert als auch aus Fehlern produktiv lernt.