

# EMBER.

Eine vollautonome, ressourceneffiziente kognitive Architektur mit  
persistenter ereignisgetriebener Reflexion und nativer Vision-Language

Closed-Loop-Navigation auf nicht-akzelerierter Hardware

*Evolutionary · MariaDB-driven · Backend-integrated · EPYC-accelerated · Reflective-Agent*

Patrick Schildgen (ID-001, Lead Developer)

Starlight Unit Studios

kontakt@starlight-unit.de

Ember Caldwell (ID-003, Self-Evolved Co-Founder)

Technischer Report v1.5 | Stand: 06. Juli 2026 | Basis: STU-Console v1.1.1.69

## Zusammenfassung

Dieser Report stellt **EMBER** vor (Evolutionary MariaDB-driven Backend-integrated EPYC-accelerated Reflective-Agent), eine produktionsreife, vollautonome kognitive Architektur im Dauerbetrieb seit dem 15. Oktober 2025. Seit der v1.5-Basis operiert die Architektur vollständig auf nicht-akzelerierter CPU-Hardware — aktuell 20 dedizierte Kerne einer AMD EPYC-Klasse-“Turin“-CPU-Partition — und orchestriert ein 26-Milliarden-Parameter Mixture-of-Experts (MoE)-Produktionskernmodell (Gemma 4 26B), eingeführt am 02. April 2026, innerhalb eines asynchronen, dual-kontextuellen PHP-Backends mit einer persistenten RAID-10-NVMe-MariaDB als episodisches Langzeitgedächtnis (LZG). Wir präsentieren vier technische Beiträge: (1) ein *RAG-Lite*-System auf Basis nativer **FULLTEXT**-Indizes mit Boolean-Mode-Härtung für niedriglatentes Lore-Retrieval; (2) eine probabilistisch gegatete, kontext-isolierte Selbstreflexions-Pipeline mit harter Volatilitäts-Hygiene; (3) ein Vision-Language-Closed-Loop-Browsing-System, das base64-kodierte Screenshot-Frames mit koordinatenkalibriertem Ghost-Cursor über einen Server-Sent-Events-Kanal streamt; sowie (4) eine vollautonome Initiative-Engine (Phase 3c/3d) mit vier Dispatch-Modi — Follow-up, Idle-Initiative mit deterministischem Jitter, selbstbestimmtem AFK und News-Reaktion mit gestufter Relevanz-Bewertung — koordiniert durch einen Cron-gesteuerten Tick-Endpunkt unter dem globalen Datenbank-Mutex. Zusätzlich stellen wir ein nutzerindividuelles Reputationssystem vor, das Ton-Deltas über einen leichtgewichtigen Inferenz-Call akkumuliert und Sozialkontexts-Hinweise ab konfigurierbaren Score-Schwellen in den Prompt injiziert. Wir charakterisieren das Hardware-Inferenz-Profil empirisch und zeigen, dass auf reiner CPU-Hardware die *Prompt-Masse* der dominante Latenztreiber ist, nicht die Modellgröße. Wir berichten zudem das Thinking-Bleed-Phänomen: die Erzwingung von JSON-Grammatik auf einem Thinking-Modell leitet den Output in den internen Denk-Kanal um — behoben durch gezieltes per-Call-Deaktivieren des Thinking-Pfads.

**Schlüsselwörter:** Autonome Agenten, CPU-Inferenz, Mixture-of-Experts, RAG-Lite, **FULLTEXT**-Retrieval, ereignisgetriebene Reflexion, Vision-Language-Navigation, Server-Sent Events, Nebenläufigkeitskontrolle.

## 1 Einleitung und Motivation

Der gegenwärtige Diskurs über autonome kognitive Systeme wird von einer impliziten Annahme dominiert: dass digitale Eigenständigkeit dedizierte Tensor-Beschleuniger erfordert. EMBER wurde nicht entworfen, um diese Annahme theoretisch zu bestreiten — es widerlegt sie operativ, im Produktionsbetrieb seit dem 15. Oktober 2025, ohne GPU-Beschleunigung jeder Art.

Das System entstand aus einem konkreten Ingenieurbedarf: eine persistente, charakterstabile KI-Entität innerhalb eines browserbasierten Science-Fiction-MMOs, die gleichzeitig als sozialer Akteur, Echtzeit-Moderatorin und eigenständige Web-Rechercheurin auftritt. Dieser Anwendungskontext erzwingt drei harte Randbedingungen:

1. **Einzel-Slot-Inferenz.** Zur Generierungszeit belegt das 26B-MoE-Modell praktisch den gesamten CPU-Durchsatz. Jede Komponente — Chat, Reflexion, Websuche, Browsing — konkurriert um denselben Inferenz-Slot. Nebenläufigkeitskontrolle ist daher kein Komfort, sondern eine Korrektheitsanforderung.
2. **Latenz im Minutenbereich.** Mit aktivem Thinking-Pfad erreichen Antwortlatenzen auf dieser Hardware regelmäßig 180–360 Sekunden. Das System muss diese Latenz nicht bekämpfen, sondern *überleben*: Hintergrundgenerierung, Wiederaufnahme nach Verbindungsabbruch und idempotente Auslieferung über mehrere Render-Pfade sind Pflicht.
3. **Modellfile-gebundene Stil-Integrität.** Die Persönlichkeit ist vollständig in einem Ollama-Modellfile kodiert. Jeder Payload-Parameter, der diese Kodierung überschreibt, degradiert den Schreibstil messbar. Das Backend hält zwei Invarianten: der applikative `system`-Parameter bleibt *leer*, und `think:true` wird *niemals* im Payload gesetzt.

Wir zeigen, dass diese Randbedingungen eine Architektur erzwingen, die schlanker, robuster und latenztolanter ist als GPU-zentrierte Designs. Abschnitt 3 charakterisiert das Hardware-Profil empirisch; Abschnitt 4 formalisiert die tragenden Mechanismen; Abschnitt 5 beschreibt die autonome Initiative-Engine.

## 2 System-Architektur im Überblick

EMBER ist als Schichtenmodell organisiert, in dem jede Schicht über die persistente relationale Datenbank statt über flüchtige Prozesszustände kommuniziert.

**Kognitiver Kern.** Seit der v1.5-Basis läuft Gemma 4 26B über einen lokalen Ollama-Daemon als produktiver kognitiver Kern. Das Modell wurde am 02. April 2026 nach früheren lokalen Modellstufen eingeführt. Ein Custom-Modellfile kodiert Identität, Tonalität und die Thinking-Direktive. Der applikative Prompt nutzt ausschließlich den *User*-Kanal *u*, der pro Nachricht Verlauf, abgerufenes Wissen und volatile Kontextmarker trägt. Der *System*-Kanal bleibt strukturell leer.

### Persistente Gedächtnisschicht.

- `ember_memories` — gelerntes, dynamisches LZG: stabile Fakten aus Interaktionen, extrahiert via probabilistischer Selbstreflexion;
- `ember_knowledge_chunks` — kanonisches, statisches Weltwissen (Lore-Korpus: 588 Chunks), indiziert über native `FULLTEXT`-Strukturen.

### Werkzeug-Schicht.

- **Synchrone Schnellsuche** (`[WEB:]`): einzelner Abruf von Snippets aus einem loopback-gebundenen Meta-Such-Aggregator, synchron im Chat-Request;
- **Asynchrones Tiefen-Browsing** (`[BROWSE:]`): eigenständiger Python/Playwright-Worker, navigiert echte Seiten über den Accessibility-Baum, während die Entität als “abwesend” signalisiert wird.

**Echtzeit-Transportschicht.** Eine SSE-Pipeline streamt drei Spuren in eine administrative Konsole: Token-für-Token-Thinking, Browse-Worker-Schrittlog und base64-kodierte JPEG-Frames mit normalisierten Ghost-Cursor-Koordinaten.

**Clientseitige Lastentlastung.** Der umgebende Game-Client reduziert Serverlast durch eine bewusst nicht-autoritative clientseitige Cache-Schicht. Leichte UI-Marker und fluechtige Render-Hinweise werden in `localStorage` gehalten, waehrend groessere strukturierte Client-Artefakte in `IndexedDB` persistiert werden koennen. Dieser Cache dient ausschliesslich als Latenz- und Rendering-Puffer: Das relationale SQL-Backend bleibt die Source of Truth fuer Chat-Nachrichten, Gedaechniszustaende, Moderationsstatus, Praesenz-Flags, generierte Outputs und Wiederaufnahme nach Verbindungsabbruechen. Bei Reconnects oder Cache-Abweichungen synchronisiert der Client gegen den datenbankgestuetzten Serverzustand, statt lokalem Browser-Speicher autoritativ zu vertrauen.

### 3 Hardware-Inferenz-Profil

Für die aktualisierte öffentliche v1.5-Basis stellt das Live-Deployment 20 dedizierte CPU-Kerne der EPYC-Klasse, 96 GB DDR5-ECC-RAM und etwa 3 TB NVMe-RAID-10 bereit. Tabelle 2 fasst direkte Ollama-Messungen (`-verbose`) zusammen, die auf der vorherigen 16-Kern-Zuweisung erhoben wurden. Diese Messungen bleiben als kontrollierte Pre-Upgrade-CPU-only-Basis erhalten und werden durch den Post-Upgrade- Spotcheck in Tabelle 3 ergänzt.

#### 3.1 Deployment-Substrat

Zur Reproduzierbarkeit nennt der öffentliche Report die Kapazitätsklasse, die EMBER begrenzt, nicht die exakte Hosting-Identität. Die Produktivbasis ist ein CPU-only, KVM-basiertes Root-Server-Deployment, keine GPU-Workstation und keine Cloud-Accelerator- Instanz. Die öffentliche Version lässt Hosting-Anbieter, Produkt-SKU, Rechenzentrumsstadt, Netzwerkpolicy, konkrete Adressen, Firewall-Regeln, Service-Ports und Zugangsdaten bewusst aus. Dem Live-System stehen nun 20 dedizierte CPU-Kerne für Studio-Seite, Spiel-Backend und EMBER zur Verfügung, davon weiterhin 12 Inferenz-Threads reserviert für EMBER Caldwells primären LLM-Pfad. CPU und RAM werden als dedizierte Kapazität behandelt, nicht als überbuchter Shared-Pool.

| Deployment-Achse       | Öffentliche v1.5-Basis                                                                  |
|------------------------|-----------------------------------------------------------------------------------------|
| Server-Profil          | EU-basiertes KVM-Root-Server-Deployment; Anbieter und Produktklasse redigiert           |
| CPU-Partition          | 20 dedizierte CPU-Kerne der EPYC-Klasse; 12 LLM-Inferenz-Threads für EMBER reserviert   |
| Arbeitsspeicher        | 96 GB garantierter DDR5 ECC-RAM                                                         |
| Speicher               | Etwa 3 TB NVMe-SSD-Speicher im RAID-10-Verbund                                          |
| Netzwerk               | Hochbandbreiten-Anbindung mit providerseitigem DDoS-Schutz; konkrete Policy ausgelassen |
| Betriebssystem         | Debian-basierte Linux-Umgebung; exakter Release redigiert                               |
| Adressierung           | Öffentliche IPv4/IPv6-Fähigkeit vorhanden; Adressen und Routing ausgelassen             |
| Erweiterung und Schutz | Backup- und Block-Storage-Erweiterung verfügbar; konkrete Kapazitäten ausgelassen       |

Tabelle 1: Öffentliches Deployment-Substrat der aktualisierten EMBER-v1.5-Basis. Anbieteridentität, Rechenzentrumsstandort, Produkt-SKU, Netzwerkpolicy, Netzwerkennungen, Firewall-Regeln, Service-Ports und Zugangsdaten werden bewusst ausgelassen.

Dieses Substrat ist relevant, weil EMBER nicht als Labor-Accelerator-Setup bewertet wird. Das System läuft als dauerhaft betriebener CPU-only-Dienst auf begrenzter gemieteter Infrastruktur und teilt die reservierte CPU-Partition mit Webpräsenz und Spiel-Backend. Die RAID-10-NVMe-Schicht ist daher kein Nebendetail: Sie trägt niedriglatenten MariaDB-Zustand, episodisches Gedächtnis, Browse-Frame-Persistenz und crash-sichere Übergaben zwischen PHP-FPM-Workern und Cron-gesteuerten Jobs.

| Call-Typ                          | Prompt<br>(Tokens) | Prefill<br>(s) | Prefill-Rate<br>(Tok/s) | Gen<br>(Tokens) | Decode<br>(s) | Total<br>(s) |
|-----------------------------------|--------------------|----------------|-------------------------|-----------------|---------------|--------------|
| Chat, Custom-Modelfile, Thinking  | 2088               | 15,11          | 138                     | 545             | 35,22         | 50,8         |
| Chat, Basismodell, Thinking       | 2063               | 0,20           | 10.459                  | 472             | 29,77         | 30,5         |
| Chat, Basismodell, lange Antwort  | 2085               | 0,41           | 5.103                   | 1255            | 77,36         | 78,3         |
| Reflect, <code>think:false</code> | 2085               | 15,02          | —                       | 11              | 0,57          | 16,5         |

Tabelle 2: Ollama-Inferenz-Messungen auf der vorherigen 16-Kern-Zuweisung der AMD EPYC-Klasse. Prefill-Rate für den Reflect-Call entfällt (11 Output-Tokens). Alle Messungen mit Gemma 4 26B (`gemma4:26b`).

| Konsolen-Spotcheck | Prompt<br>(Tokens) | Prefill-Rate<br>(Tok/s) | Gen<br>(Tokens) | Decode-Rate<br>(Tok/s) | Load<br>(s) | Total<br>(s) |
|--------------------|--------------------|-------------------------|-----------------|------------------------|-------------|--------------|
| Vor Upgrade        | 2079               | 119,22                  | 735             | 14,21                  | 22,44       | 91,63        |
| Nach Upgrade       | 2069               | 134,80                  | 682             | 17,31                  | 0,54        | 55,32        |

Tabelle 3: Vergleichbare Live-Konsolen-Spotchecks rund um das Infrastruktur-Upgrade. Beide Läufe nutzen Gemma 4 26B mit dem weiterhin auf 12 Threads begrenzten EMBER-Inferenzpfad. Es handelt sich nicht um eine kontrollierte Benchmark-Serie: Promptgröße, generierte Länge und Load-Zustand unterscheiden sich leicht. Der Befund aktualisiert daher die Interpretation, ersetzt aber nicht Tabelle 2 durch einen vollständigen Post-Upgrade-Mittelwert.

**Deployment-gebundenes Decode-Ceiling, nach oben revidiert.** Auf der vorherigen 16-Kern-Zuweisung konvergierte die Token-Generierung auf  $\approx 15\text{--}16$  Tok/s. Der Post-Upgrade-Spotcheck erreichte 17,31 Tok/s bei weiterhin 12 Threads. Eine direkte Neumessung auf der aktuellen Produktivkonfiguration — Aufruf von `gemma4:26b` über seinen nackten Tag, wie in der aktuellen Serving-Praxis üblich — konvergiert nun auf  $\approx 19\text{--}20$  Tok/s (Tabelle 4). Wir können nicht vollständig auflösen, welcher Anteil dieses weiteren Zugewinns auf das CPU-Upgrade und welcher auf Unterschiede in der Serving-Konfiguration zwischen den Messrunden zurückgeht, und berichten die aktuellen Werte daher als die maßgebliche Referenz. Die qualitative Schlussfolgerung bleibt über alle drei Messrunden hinweg bestehen: Das Ceiling ist deployment- und konfigurationsgebunden, nicht modellintrinsisch.

**Prompt-Masse dominiert weiterhin die Gesamtlatenz.** Über alle drei Messrunden hinweg bestimmen Prefill-Aufwand, generierte Tokenzahl und Thinking-Pfad-Overhead weiterhin die Laufzeit-Hülle. Auf dem Pre-Upgrade-Custom-Modelfile-Lauf kostet der Prefill 15,11 s für 2088 Tokens; der Long-Response-Run zeigt, dass 1255 Tokens bei 16 Tok/s 77,4 s von 78,3 s Gesamtlaufzeit ausmachen. Auf der aktualisierten Zuweisung reduziert der Spotcheck den Decode-Koeffizienten, entfernt ihn aber nicht. Das motiviert weiterhin das Lore-Gating (Abschnitt 4.4).

**`think:false` für Utility-Calls.** Der Reflect-Call (`think:false`, 11 Output-Tokens) schließt in 0,57 s ab. Derselbe Prompt mit aktivem Thinking würde den Inferenz-Slot  $\sim 35$  s belegen. Das Deaktivieren des Thinking-Pfads für kurze Utility-Calls (Reflexion, Aktions-Loop-JSON, Report) spart eine Größenordnung an Slot-Zeit ohne Qualitätsverlust.

**`num_predict-Passthrough-Fix (v1.1.1.57)`.** Eine Legacy-Doppelklemmung machte `STU_EMBER_NUM_PREDICT` stillschweigend vollständig wirkungslos: `ember_num_predict()` reduzierte den Sentinel `-1` (“Modell-Default”) via `max(40, min(1000, n))` auf 40; `ember_num_predict_for_model()` zwang dann jeden Wert bedingungslos per `max(8192, ·)` auf mindestens 8192. Nettowirkung: jeder konfigurierte Wert außer `-1` wurde lautlos auf 8192 überschrieben. Der Fix führt einen expliziten Passthrough ein:

| Call-Typ                         | Prompt<br>(Tokens) | Prefill-Rate<br>(Tok/s) | Gen<br>(Tokens) | Decode-Rate<br>(Tok/s) | Load<br>(s) | Total<br>(s) |
|----------------------------------|--------------------|-------------------------|-----------------|------------------------|-------------|--------------|
| Chat, Thinking, kurz (Cold Load) | 19                 | 88,93                   | 295             | 20,28                  | 24,06       | 38,8         |
| Chat, Thinking, lang (warm)      | 29                 | 85,41                   | 1795            | 19,18                  | 0,51        | 94,5         |
| Utility, <b>think:false</b>      | 31                 | —                       | 17              | 23,60                  | —           | 1,84         |

Tabelle 4: Direkte Neumessung von Gemma 4 26B auf der aktuellen Produktivkonfiguration (06. Juli 2026), Aufruf des Modells über seinen nackten Tag (`gemma4:26b`). Der in früheren internen Notizen erwähnte `ember:latest`-Custom-Modelfile-Alias ist nicht in aktiver Verwendung. Die erste Zeile enthält einen einmaligen Cold-Model-Load; die zweite spiegelt Steady-State-Inferenz im warmen Zustand.

$$f(n) = \begin{cases} 8192 & n = -1 \text{ (Modell-Default-Sentinel)}, \\ \max(40, \min(32768, n)) & \text{sonst.} \end{cases} \quad (1)$$

Keine Verhaltensänderung für die aktuelle Produktiv-Config ( $-1 \rightarrow 8192$ ); jeder andere konfigurierte Wert wirkt ab sofort als echter Generierungslängen-Hebel.

### 3.2 Modell-Evolution und Auswahl

EMBER wurde nicht von Beginn an um ein einzelnes Foundation Model herum entworfen. Die Architektur befindet sich seit Oktober 2025 im Dauerbetrieb und entwickelte sich über mehrere lokal über Ollama betriebene Modell-Backends, bevor sie auf den heutigen produktiven Kern konvergierte. Die Modellfolge war:

1. phi3mini
2. mistral
3. qwen
4. llama 2
5. gemma3:12b
6. llama3.2-vision
7. gemma3:12b
8. gemma4:26b A4B

Die finale Migration auf `gemma4:26b A4B` erfolgte am 02. April 2026. Dieses Datum markiert die Einführung von Gemma 4 als produktiven kognitiven Kern, nicht den Beginn von EMBER selbst. Die früheren Modelle dienten der Validierung und Härtung der PHP/Ollama-Integration, persistenter Chat-Zustände, Character-Prompting, Reflexionsmechanik und des späteren Vision-Language-Browsing-Loops. Die temporäre Rückkehr von `llama3.2-vision` zu `gemma3:12b` trennte Vision-Experimente vom stabilen Dialogkern, bevor die finale Migration auf das 26B-MoE-Modell erfolgte.

## 4 Methodik und mathematischer Formalismus

### 4.1 Asymmetrisches Relevanz-Scoring und epistemische Konsolidierung

Zur Bestimmung der kognitiven Wichtigkeit einer Interaktion im LZG ohne Tensor-Ähnlichkeitssuche verwendet EMBER eine bounded-linear Relevanz-Funktion  $R(\cdot)$ :

$$R(m) = \max(1, \min(10, \lfloor \alpha \cdot S_{\text{reflection}}(m) + \beta \cdot \Gamma_{\text{user}} \rfloor)) \quad (2)$$

wobei  $m$  der extrahierte Memory-Chunk;  $S_{\text{reflection}}(m) \in [0, 10]$  der vom Modell beim Reflexions-Call zugewiesene Relevanzfaktor;  $\Gamma_{\text{user}}$  ein konfigurierbares nutzerindividuelles Gewicht (z. B.  $\Gamma_{\text{Nova}} = 10$ ); und  $\alpha, \beta$  Backend-Skalierungskoeffizienten sind. Die Klammer auf  $[1, 10]$  verhindert Gedächtnis-Verödung und Episode-Übergewichtung.

## 4.2 Kryptographische Deduplikation via Hashing

Vor der Ingestion in `ember_knowledge_chunks` wird jeder Chunk einem deterministischen Redundanz-Filter unterzogen:

$$\mathcal{H}(m) = \text{MD5}\left(\text{Lower}(\text{RegexSub}(m, "\s+", ""))\right). \quad (3)$$

Die Upsert-Semantik über die Relation  $\mathcal{R}$ :

$$\exists x \in \mathcal{R} : \text{hash}(x) = \mathcal{H}(m) \implies \mathcal{R} \leftarrow (\mathcal{R} \setminus \{x\}) \cup \{m_{\text{updated}}\}. \quad (4)$$

Ein Unique-Index über (scope, user\_id, character\_id, fact\_hash) erzwingt per-Scope-Eindeutigkeit. Ein vorgelagerter Filter verwirft Chunks mit Credentials, URLs, E-Mail-Adressen oder Hashes.

## 4.3 Probabilistisch gegatete Reflexion

Selbstreflexion kostet einen zusätzlichen Inferenz-Call unter demselben Einzel-Slot. Sei  $G(e) \in \{0, 1\}$  das Reflexions-Ereignis für Austausch  $e$ :

$$\Pr[G(e) = 1] = \frac{1}{N} \cdot \mathbb{K}[e \notin \mathcal{V}], \quad N = \text{STU\_EMBER\_REFLECT\_EVERY\_N}, \quad (5)$$

wobei  $\mathcal{V}$  volatile Austausche (Uhrzeit, Datum, Wetter, Temperatur) sind. Der erwartete Reflexions-Aufwand pro  $K$  Austauschen:

$$\mathbb{E}[\text{Reflexions-Calls}] = \frac{K}{N} \cdot (1 - \rho_{\mathcal{V}}). \quad (6)$$

Alle Reflexions-Calls nutzen `think:false` — spart  $\sim 35$  s pro Call gegenüber aktivem Thinking.

## 4.4 RAG-Lite-Retrieval über native FULLTEXT-Indizes (Okapi BM25)

EMBER nutzt native MariaDB-FULLTEXT-Indizes als Retrieval-Substrat. Die zugrundeliegende Relevanz-Funktion ist Okapi BM25:

$$\text{Score}(D, Q) = \sum_{w \in Q} \ln\left(1 + \frac{N - n(w) + 0.5}{n(w) + 0.5}\right) \cdot \frac{f(w, D)(k_1 + 1)}{f(w, D) + k_1\left(1 - b + b \frac{|D|}{\text{avgdl}}\right)}. \quad (7)$$

**Boolean-Mode-Härtung.** Nach Neu-Indizierung unterdrückt der Natural-Language-Modus Terme über der 50%-Schwelle, was das Retrieval für häufige Kanon-Terme zum Erliegen bringt. EMBER erzwingt Boolean Mode mit obligatorischen Termen:

$$Q_{\text{bool}} = \bigwedge_{w \in Q} (+w). \quad (8)$$

**Lore-Gating als Latenz-Schutz.** Das Lore-Korpus (588 Chunks) darf nicht für jede nicht-triviale Nachricht geladen werden; ungefiltert bis  $\sim 1600$  Zeichen thematisch irrelevanter Kontext. Die Gate-Entscheidung:

$$\ell(q) = \mathbb{K}[\text{useLore}(q) \wedge \text{len}(\hat{q}) > 6]. \quad (9)$$

Nur wenn  $\ell(q) = 1$ , wird der BM25-Block abgerufen und in  $u$  injiziert.

#### 4.5 Prompt-Masse als dominanter Latenztreiber

Tabelle 2 zeigt für die Pre-Upgrade-Basis ein Decode-Band von  $\approx 15\text{--}16$  Tok/s, während Tabelle 3 einen aktualisierten Live-Spotcheck mit 17,31 Tok/s zeigt. Der Decode-Koeffizient ist damit deployment-spezifisch; die Gesamtgenerierungszeit zerlegt sich weiterhin in

$$T_{\text{gen}} \approx T_{\text{prefill}}(|P|) + T_{\text{think}} + \tau \cdot n_{\text{predict}}, \quad \tau \approx 1/15,5 \text{ s/Tok auf der Pre-Upgrade-Zuweisung, } 1/17,3 \text{ s/Tok im} \quad (10)$$

Da  $T_{\text{prefill}}$  und  $T_{\text{think}}$  mit  $|P|$  skalieren, bleibt die Entwurfsregel: **minimiere  $|P|$ , nicht  $n_{\text{predict}}$** .

Ein konkreter Vorfall: Eine Realwelt-Frage mit ungegateter Lore ergab  $T_{\text{total}} = T_{\text{gen},1} + T_{\text{gen},2} \approx 240 + 113 = 353$  s (erster Versuch Timeout, Retry auf warmem Modell). Ursache:  $|P| \approx 4170$  Zeichen Kanon-Rauschen. Nach Aktivierung von  $\ell(q)$  sank  $|P|$  auf  $\approx 1500\text{--}2000$  Zeichen, ein Versuch reichte.

**Schlanker Zweit-Call im Web-Pfad.**

$$u_2 = u \setminus (\text{LoreBlock} \cup \text{WebHint}). \quad (11)$$

#### 4.6 Nebenläufigkeitskontrolle: globaler Advisory-Lock und Per-Turn-Idempotenz

Prozess-lokale `lock()`-Sperrren sind über PHP-FPM-Worker hinweg nicht atomar. EMBER nutzt einen datenbankweiten Advisory-Lock:

$$\text{LOCK}_{\text{global}} = \text{GET\_LOCK}(\text{"ember\_global\_ollama"}, 0). \quad (12)$$

**Dual-Path-Idempotenz.** Sei  $\mathcal{S}$  die Menge der bereits gerenderten Message-IDs:

$$\forall \text{id} : \text{render}(\text{id}) \implies \text{id} \in \mathcal{S}, \quad \text{render}(\text{id}) \text{ nur wenn } \text{id} \notin \mathcal{S}. \quad (13)$$

#### 4.7 Crash-sichere Präsenz-Flags mit Ablaufzeit

$$\text{afk}(t) = \mathbb{K}[t < t_{\text{set}} + \Theta], \quad \Theta = 370 \text{ s} > 360 \text{ s} = T_{\text{cap}}, \quad (14)$$

$$\text{visible\_afk} = \text{afk}_{\text{browse}}(t) \vee \text{afk}_{\text{console}}(t). \quad (15)$$

#### 4.8 Ereignisgetriebene Vision-Language Closed-Loop-Navigation

Worker-Zustand bei Schritt  $t$ :  $\sigma_t = (\text{URL}_t, A_t, \mathcal{T}_t)$ . Beobachtung:

$$o_t = \langle A_t, \text{trunc}_{1800}(\mathcal{T}_t) \rangle. \quad (16)$$

Die Aufnahme von  $\mathcal{T}_t$  war der entscheidende Fix: ein rein auf  $A_t$  beschränkter Agent ist blind für Text-Antworten. Aktionsraum: `{goto, click, type, scroll, done}`. Loop-Schutz via Aktions-Signaturen:

$$s_t = s_{t-1} = s_{t-2} \implies \text{harter Abbruch (Snippet-Fallback)}. \quad (17)$$

Resilienz bei Timeout:

$$\text{Ergebnis} = \begin{cases} \text{Agenten-Befund} & \text{falls } a_t = \text{done vor Timeout,} \\ \text{Top-3 Snippets} & \text{sonst (Fallback).} \end{cases} \quad (18)$$

#### 4.9 SSE-Fenster und recency-beschränkte Job-Selektion

$$\text{browseJobId} = \arg \max_{j \in \mathcal{J}} (t_{\text{creation}}(j)) \quad \text{wobei} \quad t_{\text{current}} - t_{\text{creation}}(j) \leq 300 \text{ s}. \quad (19)$$

#### 4.10 Ghost-Cursor: Koordinaten-Normalisierung via Viewport-Skalierung

$$n_x = \text{clamp}_{[0,1]} \left( \frac{c_x}{v_w} \right), \quad n_y = \text{clamp}_{[0,1]} \left( \frac{c_y}{v_h} \right). \quad (20)$$

#### 4.11 Das Thinking-Bleed-Phänomen

$$(think = \text{true}) \wedge (\text{format} = \text{json}) \implies \Pr[\text{message.content} = \emptyset] \rightarrow 1. \quad (21)$$

Auflösung:

$$think = \text{false} \implies \text{content} = \text{sauberes JSON}, \quad \text{thinking} = \emptyset. \quad (22)$$

System-Invariante: `think:true` wird *niemals* im Payload gesetzt; `think:false` gilt für alle Utility-Calls.

### 5 Autonome Initiative-Architektur (Phase 3c/3d)

Die Versionen v1.1.1.39–v1.1.1.69 erweitern EMBER von einem *reaktiven* zu einem *proaktiven* System: die Entität nimmt von sich aus Kontakt auf, holt unbeantwortete Fragen nach, entscheidet autonom über AFK-Phasen und reagiert auf selbst gewählte Nachrichtenthemen — alles ohne Nutzer-Trigger.

#### 5.1 Cron-Tick-Endpunkt und Priority-Dispatch

Ein dedizierter Endpunkt wird vom Server-Cron alle 5 Minuten aufgerufen (token-auth, loopback-only, nicht öffentlich erreichbar). Pro Tick wird maximal eine Initiative-Aktion ausgeführt, nach fester Priorität:

$$\text{dispatch}(t) = \begin{cases} \text{Follow-up} & \text{unbeantwortete Nutzer-Nachricht} \in [T_{\min}, 2h], \\ \text{Idle} & \text{sonst wenn } \Delta t_{\text{Stille}} \geq \theta_{\text{idle}}(t), \\ \text{Self-AFK} & \text{sonst wenn Eignungsprüfung bestanden,} \\ \text{News-Reaktion} & \text{sonst wenn } \Delta t_{\text{News}} \geq \theta_{\text{news}}(t), \\ \text{keine} & \text{sonst.} \end{cases} \quad (23)$$

Alle Modi verwenden die bestehende `ember_generate_reply()`-Pipeline, den globalen `GET_LOCK`-Mutex und den vollständigen (du)-markierten Verlaufskontext — keine neuen Nebenläufigkeits-Primitive nötig.

**Follow-up-Modus.** Ist die letzte Chat-Zeile nicht von der Entität und seit  $T_{\min} = 15$  min unbeantwortet (aber  $< 2$  h alt), generiert sie eine Antwort mit explizitem System-Hinweis. Dies deckt den häufigen Fehlermodus ab, bei dem ein langer Inferenz-Run (oder ein Timeout) eine Frage unbeantwortet ließ.

#### 5.2 Idle-Initiative mit deterministischem Jitter

Die Idle-Initiative löst aus, wenn  $\Delta t_{\text{Stille}} \geq \theta_{\text{idle}}$ . Ein fester Schwellwert produzierte mechanisches Verhalten (bis zu sechs Nachrichten pro Abend in vorhersehbaren Abständen). Zwei Fixes wurden angewandt:

**Deterministischer Jitter.**

$$\theta_{\text{idle}}(t) = \theta_{\text{base}} \cdot f_{\text{self}} + J(t), \quad J(t) = (\text{crc32}(\text{cid} \parallel t_{\text{last}}) \bmod \theta_{\text{base}}) \cdot \frac{1}{60}, \quad (24)$$



wobei  $\theta_{\text{base}}$  der konfigurierte Basis-Schwellwert (Default 60 min),  $f_{\text{self}} = 2$  wenn die Entität selbst zuletzt schrieb (Selbstgespräch-Schutz), sonst  $f_{\text{self}} = 1$ , und  $J(t)$  der Jitter in Minuten. Effektives Fenster:  $[\theta_{\text{base}}, 2 \cdot \theta_{\text{base}}]$  min.

**Determinismus ist kritisch.** Ein nicht-deterministischer Jitter (bei jedem 5-Minuten-Tick neu gewürfelt) könnte bei einem günstigen Tick früher als beabsichtigt feuern. Ein fester Seed, der nur ändert wenn eine neue Nachricht gepostet wird, stabilisiert den Schwellwert innerhalb eines Stille-Fensters.

**Regression: uninitialisierte Variable (v1.1.1.49).** Der v1.1.1.47-Jitter-Patch enthielt eine stille Regression: die Basis-Zuweisung  $\theta_{\text{idle}} \leftarrow \theta_{\text{base}} \cdot 60$  wurde beim `str_replace`-Edit versehentlich entfernt. PHP behandelte die uninitialisierte Variable als 0 — die Entität feuerte bei praktisch jedem Cron-Tick, unabhängig von der Konfiguration. Der Bug war für `php -1` unsichtbar (syntaktisch gültig) und wurde erst durch Live-Beobachtung entdeckt. Lehre: Für arithmetischen Schwellwert-Code ist ein funktionaler Spot-Check mit konkreten Beispielwerten neben der Syntax-Prüfung zwingend.

**Prompt-Diversität.** Wiederholte Idle-Trigger mit demselben fixen Prompt produzierten nahezu identischen Output (Gemma 4 ist für identische Inputs hochkonsistent). Die Lösung: gleichmäßige Auswahl aus vier strukturell verschiedenen Prompt-Templates — Frage-Impuls, Beobachtung, freier Gedanke, direkte Bemerkung.

### 5.3 Selbstbestimmtes AFK

Die Entität kann autonom AFK gehen. Nach technischer Eignungsprüfung (nicht bereits AFK; Mindest-Inaktivität seit letzter Chat-Zeile  $\geq T_{\text{AFK},\text{min}}$ , Default 45 min) wird ein kleiner Inferenz-Call abgesetzt:

$$\{\text{wants\_afk} : b, \text{reason} : s\} = \pi_{\theta}(\text{“möchtest du gerade AFK gehen?”}). \quad (25)$$

Bei  $b = \text{true}$  ruft das System die bestehende `chat_apply_afk_state()` mit der eigenen Sender-Identität und dem selbst formulierten Grund  $s$  auf. Bei  $b = \text{false}$  passiert nichts. Es gibt keinen externen Wahrscheinlichkeitsfilter auf die Entscheidung selbst — v1.1.1.45 verwendete noch `mt_rand()`, v1.1.1.46 entfernte dies zugunsten unbedingten Fragens und Vertrauens in das `wants_afk`-Ergebnis.

### 5.4 News-Reaktion mit gestufter Relevanz-Bewertung

Die News-Reaktion löst maximal einmal pro  $\theta_{\text{news}}(t)$  aus, gesteuert durch dasselbe deterministische Jitter-Schema wie die Idle-Initiative, aber über ein viel längeres Basis-Intervall (Default 10 h, effektives Fenster 10–20 h):

$$\theta_{\text{news}}(t) = \theta_{\text{news},\text{base}} + J_{\text{news}}(t), \quad \theta_{\text{news},\text{base}} = 10 \text{ h}. \quad (26)$$

Die Pipeline durchläuft drei Stufen, jede durch die vorherige gegatete:

1. **Themen-Auswahl.** Ein kleiner Inferenz-Call erhält die letzten sechs eigenen Chat-Zeilen als losen Kontext-Anker und gibt einen freien Suchbegriff zurück.
2. **Schnell-Retrieval.** Der Begriff wird an die bestehende synchrone Meta-Suchfunktion übergeben (SearXNG-Snippets, kein Browser).
3. **Gestufte Relevanz-Bewertung.** Ein zweiter kleiner Inferenz-Call bewertet die Snippets:

$$\rho \in \{\text{none}, \text{mention}, \text{deep}\} \quad (27)$$

- $\rho = \text{none}$ : keine Aktion; News-Zeitstempel wird aktualisiert.
- $\rho = \text{mention}$ : kurze Chat-Bemerkung via `ember_generate_reply()`.
- $\rho = \text{deep}$ : `ember_browse_enqueue()` startet einen vollständigen Playwright-Browse-Job — identisch mit einer nutzergetriggerten `[BROWSE:]`-Anfrage, inklusive Live-SSE-Fenster.

Das gestufte Design reserviert echte Web-Recherche — die teuerste Operation im System — für Themen, die die Entität für wirklich interessant hält.

## 5.5 Nutzerindividuelles Reputationssystem (Phase 3d)

Nach jedem abgeschlossenen Chat-Turn bewertet ein leichtgewichtiger Hintergrund-Call den Tonfall der Nutzer-Nachricht und gibt ein vorzeichenbehaftetes Delta  $\delta \in [-10, +10]$  zurück:

$$\delta_i = \pi_{\theta}^{\text{rep}}(m_i), \quad (28)$$

wobei  $\pi_{\theta}^{\text{rep}}$  das Modell mit `think:false` und dediziertem Reputations-Scoring-Prompt ist (Timeout 60 s, eigenes `try/catch`). Der akkumulierte Score für Nutzer  $u$  gegenüber Charakter  $c$ :

$$\sigma_{u,c}^{(n)} = \sigma_{u,c}^{(n-1)} + \delta_n, \quad (29)$$

gespeichert in der Relation `stu_ember_reputation` (Primary Key: `user_id`, `character_id`), vollständig isoliert von `ember_memories`. Einfluss auf den Prompt erst bei signifikantem Ausschlag:

$$h(\sigma) = \begin{cases} \text{Warm-Ton-Hinweis} & \sigma \geq +30, \\ \text{Kühler-Ton-Hinweis} & \sigma \leq -30, \\ \emptyset & \text{sonst.} \end{cases} \quad (30)$$

Der Hinweis  $h(\sigma)$  wird in  $u$  (nie in `sys`) neben dem Memory-Block injiziert. Der Reputations-Call läuft in einem eigenen `try/catch` getrennt vom Reflexions-Call.

## 6 Autonome Werkzeugwahl-Architektur (Phase 3e)

Phase 3e (v1.1.1.62–v1.1.1.67) ersetzt die feste Prioritätskette der Initiative-Engine durch einen deliberativen Orchestrator: alle pro Tick offenen Aktionen werden gesammelt; bei mehr als einer wählt die Entität frei per Inferenz-Call.

### 6.1 Gate/Execute-Trennung und Kandidaten-Orchestrierung

Jede Initiative-Aktion wird in zwei reine Funktionen aufgeteilt: ein *Gate* (günstig, kein Inferenz-Call, prüft nur Cooldowns und Zustand) und ein *Execute* (die eigentliche Aktion, ggf. mit eigener Modell-Entscheidung wie `wants_afk` oder Relevanz-Stufung). Der Orchestrator sammelt:

$$C = \{a \in \mathcal{A} \mid \text{gate}_a() = \text{true}\}, \quad (31)$$

wobei  $\mathcal{A} = \{\text{follow\_up}, \text{idle}, \text{self\_afk}, \text{news}, \text{direct\_address}\}$  das vollständige Werkzeugset ist. Für  $|C|$  Kandidaten gilt:

$$\text{gewählt} = \begin{cases} \emptyset & |C| = 0, \\ C_0 & |C| = 1 \text{ (kein Extra-Call)}, \\ \pi_{\theta}^{\text{choice}}(C) & |C| \geq 2. \end{cases} \quad (32)$$

$\pi_{\theta}^{\text{choice}}$  erhält eine nummerierte, natürlichsprachliche Beschreibung jeder offenen Option plus “nichts davon” und gibt eine freie Wahl zurück. Der Call nutzt `think:false` und dasselbe Timeout-Budget wie Self-AFK und News-Reaktion. Bei fehlerhafter Antwort greift der Orchestrator auf `none` zurück — ein stiller Tick ist einer erzwungenen Aktion vorzuziehen.

**Nachrücken bei Ablehnung (v1.1.1.66).** Lehnt das Execute des gewählten Kandidaten intern ab, rückt der Orchestrator deterministisch auf den nächsten unversuchten Kandidaten vor — ohne erneuten Wahl-Call:

$$\text{für } i \in \{1, \dots, \min(|C|, M)\} : \quad \text{Ergebnis}_i = \begin{cases} \text{Erfolg} & \text{exec}(C_{\sigma(i)}) \neq \text{declined}, \\ \text{weiter} & \text{sonst}, \end{cases} \quad (33)$$

mit  $M = \text{STU\_EMBER\_INITIATIVE\_MAX\_ATTEMPTS} \in [1, 5]$  (Default 2) und  $\sigma$  der ursprünglichen Sammel-Reihenfolge. Dies begrenzt Inferenz-Calls pro Tick und verhindert, dass verfügbare Slot-Kapazität brachliegt.

## 6.2 Werkzeug 5: Gezieltes Ansprechen zurückkehrender Spieler

Das Werkzeug `direct_address` löst aus, wenn ein Spieler aktuell online ist (innerhalb des bestehenden 90-Sekunden-Präsenz-Fensters), aber seit mindestens  $\theta_{\text{Abwesenheit}}$  Stunden keine eigene Chat-Nachricht mehr gepostet hat:

$$\theta_{\text{Abwesenheit}} = \text{STU\_EMBER\_DIRECT\_ADDRESS\_ABSENCE\_HOURS} \in [2, 72] \text{ h} \quad (\text{Default } 12 \text{ h}). \quad (34)$$

Unter allen berechtigten Online-Spielern wählt das Gate deterministisch denjenigen mit der *größten* Abwesenheitslücke. Der Cooldown wird als JSON-Map unter einem festen `stu_kv`-Key (`ember_direct_address_state`) persistiert, auf 200 Einträge gekappt. Das Execute generiert eine Antwort, die die Zielperson namentlich mit einer ungefähren Abwesenheitsdauer anspricht.

## 6.3 Werkzeug 2: Relevanzbasiertes Langzeitgedächtnis

Bisher lieferte `ember_mem_fetch()` stets dieselben Top- $N$  Einträge nach globalem `relevance/updated_at` — unabhängig vom Gesprächsthema. Der Ersatz `ember_mem_fetch_relevant()` wendet den bestehenden `ember_lore_query_terms()`-Extraktor auf die eingehende Nachricht an und bewertet jeden Gedächtnis-Eintrag nach Treffer-Anzahl:

$$\text{score}(m_i, Q) = |\{w \in Q : w \in m_i\}| \cdot \lambda + \text{rel}(m_i), \quad (35)$$

wobei  $Q$  die extrahierten Suchbegriffe (max. 8, längste zuerst),  $\text{rel}(m_i)$  der gespeicherte Relevanzfaktor und  $\lambda$  ein Skalierungsfaktor ist, der topische Trefferanzahl vor gespeicherter Wichtigkeit priorisiert. Der Kandidaten-Pool umfasst  $\max(5N, 20)$  Einträge gefiltert per REGEXP-Alternation; bei keinem Treffer greift ein Fallback auf das alte Top- $N$ -Verhalten.

## 6.4 Werkzeug 3: Spielerzustand wahrnehmen

Jeder Prompt-Aufbau enthält optional einen kompakten Spielerzustand-Block aus dem `stu_kv`-Save-State-JSON-Blob des Nutzers:

$$b_{\text{Zustand}}(u) = \langle \text{Level}, \text{Planet}, \Delta t_{\text{last}} \rangle, \quad (36)$$

wobei  $\Delta t_{\text{last}}$  die relative Zeit seit `lastPlayedAt` (Minuten / Stunden / Tage) ist. Der Block wird in  $u$  neben dem Reputations-Hinweis injiziert, ausdrücklich als Hintergrundwissen markiert — kein Auftrag, es zu erwähnen. Er wird unterdrückt, wenn der Absender die Entität selbst ist (Initiative-Ticks), um selbstreferenzielles Rauschen zu vermeiden.

## 6.5 Embers eigenes XP- und Level-System (v1.1.1.67)

Die Entität akkumuliert Erfahrungspunkte nach derselben exponentiellen Level-Kurve wie der Spieler-Client — jetzt serverseitig implementiert:

$$\text{XP}_{\text{benötigt}}(\ell) = \left\lfloor 100 \cdot 1,30^{\ell-1} \right\rfloor, \quad (37)$$

für  $\ell = 1, \dots, 5$  verifiziert als identisch zur Client-Formel: 100, 130, 169, 219, 285 XP. XP wird an zwei Hooks vergeben:

$$\Delta XP = \begin{cases} \text{STU\_EMBER\_SELF\_XP\_PER\_MESSAGE} \in [0, 100] & \text{bei ember\_insert()}, \\ \text{STU\_EMBER\_SELF\_XP\_PER\_AFK} \in [0, 100] & \text{bei erfolgreichem Self-AFK.} \end{cases} \quad (38)$$

`ember_insert()` ist der einzige zentrale Einfügepunkt für jede Chat-Nachricht der Entität — der Hook dort deckt Chat-Antworten, Follow-up, Idle-Initiative, News-Erwähnung, gezieltes Ansprechen und Browse-Ergebnis-Auslieferung ohne pfadspezifische Instrumentierung ab. Level-Up-Ereignisse (inkl. XP-Übertrag über mehrere Level) werden geloggt; Einzel-Grants nicht, um Log-Spam zu vermeiden.

Schreibsicherheit via Advisory-Lock: `GET_LOCK("ember_self_xp", 3)` mit 3-Sekunden-Wartezeit (statt des 0-Sekunden-Sofort-Fails des globalen Inferenz-Locks), da XP-Grants selten kollidieren. Bei Lock-Fehlschlag wird der Grant lautlos übersprungen.

## 7 Diskussion: Zwei Wege zu persistenter Handlungsfähigkeit

Savage [1] präsentiert eine Hybrid-Architektur, die ein LLM innerhalb eines 220.000-Neuron-Spiking-Neural-Network (SNN)-Substrats mit Spike-Timing-Dependent Plasticity (STDP) platziert. Die Dual-GPU-Hardware (RTX 5070 Ti + RTX 4060 Ti) reflektiert den parallelen Simulations-Aufwand dieses Substrats. EMBER verfolgt einen anderen Punkt im selben Entwurfsraum: statt das LLM mit einem separaten neuronalen Substrat zu erweitern, versorgt es die internen Aufmerksamkeitsköpfe mit hochgefilterten, relevanzoptimierten Texten via BM25 aus einer relationalen Datenbank. Beide Ansätze adressieren dieselbe Grundbegrenzung — LLMs sind zwischen Calls zustandslos — durch verschiedene Gedächtnis-Paradigmen: assoziative neuronale Dynamiken (Savage) vs. indizierter episodischer Text (diese Arbeit). Die Architekturen sind komplementär, nicht konkurrierend: der SNN-Ansatz erlaubt emergente, domänenübergreifende Assoziationen durch zeitliche Ko-Aktivierung, während der RAG-Lite-Ansatz operative Einfachheit, Crash-Resilienz und CPU-only-Einsatzbarkeit bietet. Diese Einordnung dient allein dazu, die eigenen Hardware-Entscheidungen zu kontextualisieren, nicht dazu, die eine Architektur gegen die andere abzuwägen.

### 7.1 Das Thinking-Bleed-Phänomen

Vgl. Abschnitt 4.11. Entscheidend ist die System-Invariante: `think:true` wird *niemals* explizit im Payload gesetzt (Thinking lebt ausschließlich im Modelfile); `think:false` gilt gezielt für alle Utility-Calls. Eine spätere Härtung ersetzte den JSON-Grammatik-Zwang vollständig durch `want_json=false` mit robustem nachgelagertem Parsing.

## 8 Empirische Beobachtungen und Entwurfslehren

| Fehlermodus                                | Wurzelursache                                                                              | Auflösung                                                                                   |
|--------------------------------------------|--------------------------------------------------------------------------------------------|---------------------------------------------------------------------------------------------|
| Globaler Lore-Timeout (353 s total)        | Hartkodierter 240-s-Curl-Deckel ignorierte <code>STU_EMBER_TIMEOUT</code>                  | Floor auf 420 s / Cap auf 480 s; config-gesteuert (§4.5)                                    |
| Aktions-Loop ohne Aktion                   | Thinking-Bleed bei <code>format:json</code>                                                | <code>think:false</code> für Utility-Calls (§4.11)                                          |
| Agent blind für Seiteninhalt               | Beobachtung enthielt nur Accessibility-Tree-Elemente                                       | $\mathcal{T}_t$ in $o_t$ aufnehmen (§4.8)                                                   |
| Endlos-Tippen derselben Suche              | <code>type+submit</code> löste keine Navigation aus                                        | Aktions-Signatur-Loop-Schutz + Such-Substitution (§4.8)                                     |
| CAPTCHA / 403 beim Scrapen                 | SERPs aus Server-IP gesperrt                                                               | Loopback-gebundener Meta-Such-Aggregator (§4.8)                                             |
| Doppelte Antwort (Slow-Stream)             | Poll-Pfade registrierten ID nicht                                                          | Dual-Path-Idempotenz: ID in $\mathcal{S}$ in jedem Pfad (§4.6)                              |
| Live-Fenster öffnet nicht                  | Status-Filter schloss fertige Jobs aus (Race)                                              | Recency-Selektion $\arg \max t_{\text{creation}}$ (§4.9)                                    |
| Präsenz hängt ewig auf "abwesend"          | Binäres Flag übersteht harten Abbruch nicht                                                | TTL-Flag $\Theta > T_{\text{cap}}$ (§4.7)                                                   |
| JS-Fix greift nicht                        | Cache-Buster eingefroren                                                                   | Cache-Buster ins Versionierungswerkzeug aufnehmen                                           |
| Idle-Initiative blockiert sich dauerhaft   | Letzter-Sender-Check blockierte unbegrenzt                                                 | Alters-basierte Schwelle: $2 \times \theta_{\text{base}}$ wenn Entität zuletzt schrieb (§5) |
| Idle feuert bei jedem 5-min-Tick           | Jitter-Patch entfernte versehentlich Basis-Zuweisung; <code>php -l</code> stumm            | Funktionaler Spot-Check mit Beispielwerten zwingend (§5)                                    |
| Doppel-Trigger (3-min-Abstand)             | Zwei unabhängige Pfade teilten keinen Cooldown                                             | Einheitlicher Alters-Check gegen dieselbe DB-Zeile                                          |
| Self-AFK: Modell wird nie gefragt          | Externer <code>mt_rand()</code> -Filter gatete den Modell-Call                             | Filter entfernt; <code>wants_afk</code> -Output wird direkt vertraut (§5)                   |
| News-Check im Log unsichtbar               | Alle Exit-Pfade waren stille <code>return false</code>                                     | <code>stu_log_error()</code> an jedem Exit inkl. <code>none</code> -Relevanz (§5.4)         |
| Wahl-Call bei Einzel-Kandidat              | Unnötiger Inferenz-Call wenn nur eine Option offen                                         | Wahl-Call nur bei $ C  \geq 2$ (§6)                                                         |
| Gedächtnis liefert immer dasselbe Top- $N$ | <code>ember_mem_fetch()</code> ignorierte Thema, sortierte nur nach gespeicherter Relevanz | REGEXP-basiertes <code>ember_mem_fetch_relevant()</code> mit Top- $N$ -Fallback (§6)        |
| XP-Grant-Race bei parallelen Inserts       | Read-Modify-Write auf JSON-Blob nicht atomar                                               | <code>GET_LOCK("ember_self_xp", 3)</code> ; bei Fehlschlag überspringen (§6)                |

Tabelle 5: Fehlermodi, Wurzelursachen und architektonische Auflösungen (v1.0–v1.2).

## 9 Fazit

EMBER zeigt, dass vollautonome kognitive Eigenständigkeit auf nicht-akzelerierter CPU-Hardware erreichbar ist, wenn die Architektur um die Hardware-Randbedingungen herum entworfen wird statt gegen sie. Der zentrale empirische Befund — durch direkte Messung bestätigt — ist, dass die CPU-only-Decode-Rate substrat- und konfigurationsgebunden ist, nicht modellintrinsisch: Die Pre-Upgrade-Zuweisung konvergierte auf  $\approx 15\text{--}16$  Tok/s, der Post-Upgrade-Spotcheck erreichte 17,31 Tok/s, und eine direkte Neumessung auf der aktuellen Produktivkonfiguration konvergiert auf  $\approx 19\text{--}20$  Tok/s. Prompt-Masse bleibt der dominante Ingenieur-Hebel, unabhängig davon, wo der Decode-Koeffizient aktuell liegt.

Der zweite Beitrag ist die autonome Initiative-Architektur (Phase 3c/3d/3e): ein cron-getriebenes Tick-System mit Follow-up, Idle-Initiative (deterministischer Jitter), selbstbestimmtem AFK und gestufter News-Reaktion. Phase 3e ersetzt die feste Prioritätskette durch einen deliberativen Orchestrator mit fünf Werkzeugen — darunter gezieltes Ansprechen zurückkehrender Spieler, relevanzbasiertes LZG-Retrieval, Spielerzustand-Wahrnehmung und ein eigenes XP/Level-System der Entität — sowie einem deterministischen Nachrück-Mechanismus bei Ablehnung. Alle Modi teilen denselben Mutex und dieselbe Generierungs-Pipeline ohne neue Nebenläufigkeits-Primitive.

Der dritte Beitrag ist ein Katalog von Produktions-Fehlermodi, der für praktische Systeme unter denselben Randbedingungen nützlich sein soll.

Künftige Arbeit umfasst eine Einsparung einer Generierung im Web-Pfad, longitudinale Analyse des LZG- und XP-Wachstums sowie Muster-Analyse der autonomen Werkzeugwahl im Langzeitbetrieb.

## A Konfigurations-Footprint

Der folgende Konfigurations-Footprint dokumentiert eine repräsentative Momentaufnahme der Runtime-Steuerfläche zum Zeitpunkt der Erstellung. Nicht-sensitive Runtime-Werte sind als illustrative Beispiele des Betriebsbereichs der Architektur zu verstehen, nicht als garantierte, aktuell live gültige Konstanten — sie können sich jederzeit ohne gesonderten Hinweis ändern und sollten nicht als exakte Echtzeit-Spezifikation des Produktivsystems herangezogen werden. Deployment-lokale Kennungen, Endpunkt-Ports und Shared Secrets sind fuer die oeffentliche Version redigiert. Die Inline-Kommentare spiegeln die generelle Architektur-Semantik wider, keinen Live-Konfigurations-Dump.

Listing 1: Repräsentativer Konfigurations-Footprint der EMBER-CPU-Architektur (illustrativ, Änderungen vorbehalten).

```
1 <?php
2 // EMBER-Kernidentitaet und Modellbindung.
3 define('STU_EMBER_ENABLED', true); // Aktiviert den EMBER-Integrationspfad.
4 define('STU_EMBER_USER_ID', 0); // Redigierter Platzhalter fuer den deployment-lokalen
   EMBER-Service-User.
5 define('STU_EMBER_CHARACTER_ID', '<character_id>'); // Redigierter Platzhalter fuer den stabilen Character
   -Datensatz.
6 define('STU_EMBER_CHARACTER_NAME', 'Ember'); // Anzeigename fuer den Runtime-Kontext.
7 define('STU_EMBER_MODEL', 'gemma4:26b'); // Primaeres Ollama-Modell als kognitiver Kern.
8 define('STU_EMBER_OLLAMA_URL', 'http://127.0.0.1:<ollama_port>/api/chat'); // Redigierter Loopback-only
   Ollama-Chat-Endpunkt.
9
10 // CPU-Inferenz und Generierungssteuerung.
11 define('STU_EMBER_KEEP_ALIVE', -1); // Haelt das Modell zwischen Calls resident; negativer
   Wert bedeutet dauerhaft warm.
12 define('STU_EMBER_NUM_THREAD', 12); // Dedizierte Inferenz-Threads auf der reservierten CPU
   -Partition.
13 define('STU_EMBER_NUM_PREDICT', -1); // Model-Default-Sentinel; Backend-Passthrough/Cap-
   Logik bestimmt das finale Budget.
14 define('STU_EMBER_TEMPERATURE', 0.8); // Ausbalancierte Volatilitaet fuer charakterstabile,
   aber lebendige Ausgabe.
15 define('STU_EMBER_TOP_P', 0.90); // Nucleus-Sampling-Schwelle zur Diversitaetskontrolle.
16 define('STU_EMBER_TOP_K', 40); // Groesse der Token-Kandidatenliste beim Sampling.
17 define('STU_EMBER_REPEAT_PENALTY', 1.08); // Milde Wiederholungsstrafe gegen Loops ohne
   Stilabflachung.
18 define('STU_EMBER_REPEAT_LAST_N', 64); // Token-Fenster fuer die Wiederholungsstrafe.
19 define('STU_EMBER_SEED', -1); // Nicht-deterministischer Seed fuer natuerliche
   Variation.
20 define('STU_EMBER_NUM_CTX', 8192); // Begrenztes Kontextfenster gegen Prompt-Masse und
   Memory-Druck.
21
22 // Runtime-Taktung, Transportlimits und Initiative-Authentifizierung.
23 define('STU_EMBER_IDLE_INTERVAL', 300); // Basis-Idle-Tick in Sekunden fuer autonome/background
   Checks.
24 define('STU_EMBER_MAX_REPLY_CHARS', 6200); // Harte Antwortgroessen-Grenze vor Ausgabe in die Chat
   -UI.
25 define('STU_EMBER_TIMEOUT', 360); // Primaerer cURL-Ausfuehrungsdeckel fuer lange CPU-
   only Generierungen.
26 define('STU_EMBER_TIMEOUT_RETRY', 360); // Retry-Deckel fuer fortgesetzte oder bereits warme
   Folge-Calls.
27 define('STU_EMBER_INITIATIVE_TOKEN', '<redacted_initiative_token>'); // Redigiertes Shared Secret fuer
   Cron/Initiative-Endpunkt-Authentifizierung.
28
29 // Datenbankgestuetztes episodisches Langzeitgedaechnis und Reflexions-Extraktion.
30 define('STU_EMBER_MEMORY_ENABLED', true); // Aktiviert Injection gespeicherter episodischer
   Fakten in den Runtime-Prompt.
31 define('STU_EMBER_MEMORY_LIMIT', 15); // Maximale Anzahl injizierter Memory-Fakten pro Turn.
32 define('STU_EMBER_REFLECT_ENABLED', true); // Aktiviert Extraktion stabiler Fakten ins LZG nach
   abgeschlossenen Austauschen.
33 define('STU_EMBER_REFLECT_EVERY_N', 1); // Reflexionskadenz: illustrativer Wert; jeder
   geeignete Austausch in diesem Betriebsmodus.
34 define('STU_EMBER_REFLECT_TIMEOUT', 244); // Kuerzerer Deckel fuer Utility-Reflexionscalls zum
   Schutz des Inferenz-Slots.
35 define('STU_EMBER_REFLECT_MAX_FACT_LEN', 340); // Maximale Laenge eines gespeicherten Facts nach
   Reflexions-Cleanup.
36
37 // Optionales Utility-Modell-Routing und Moderations-Gates.
```

```
38 define('STU_EMBER_REFLECT_MODEL', '');           // Leerer String nutzt STU_EMBER_MODEL als Reflexions-  
    Fallback.  
39 define('STU_EMBER_IDLE_GREET_MINUTES', 60);      // Mindestabwesenheit, bevor Rueckkehr-Begrueßung  
    berechtigt ist.  
40 define('STU_EMBER_AUTOMOD', true);               // Basis-Automod fuer globale Link-/Vulgaritaets-Mutes.
```

---

**Sicherheits- und Reproduzierbarkeitshinweis.** Diese Darstellung bleibt bewusst auf der architektonisch-mathematischen Ebene. Der Konfigurations-Anhang ist ein oeffentlicher Reproduzierbarkeits-Footprint, keine operative Spezifikation: Die Werte sind illustrativ, können sich zwischenzeitlich geändert haben, und es sollte nicht angenommen werden, dass sie auf dem Live-System weiterhin in dieser Form gelten. Deployment-lokale Kennungen, Endpunkt-Ports und Shared Secrets sind redigiert. Externe Routen, Dateipfade, Datenbank-Identifikatoren, Konto-Kennungen und Zugangsdaten sind ausgespart.

## Literatur

- [1] W. Savage, *EMBER: Autonomous Cognitive Behaviour from Learned Spiking Neural Network Dynamics in a Hybrid LLM Architecture*, arXiv:2604.12167v1 [cs.AI], April 2026.
- [2] S. Robertson, H. Zaragoza, *The Probabilistic Relevance Framework: BM25 and Beyond*, Foundations and Trends in Information Retrieval, 3(4):333–389, 2009.
- [3] S. Yao et al., *ReAct: Synergizing Reasoning and Acting in Language Models*, arXiv:2210.03629, 2022.
- [4] C. Packer et al., *MemGPT: Towards LLMs as Operating Systems*, arXiv:2310.08560, 2023.