



Holger Reibold

Synergie der Intelligenz

Das Handbuch für das Design
und die Implementierung von
Multi-Agenten-Systemen

BRAIN-MEDIA.DE

**Implementierungs-
Blueprints**

Implementierungs-Blueprints übersetzen die theoretischen Konzepte von Multi-Agenten-Systemen in belastbare technische Grundmuster. Ihr Zweck besteht darin, nicht bei null beginnen zu müssen, sondern auf bewährte Referenzarchitekturen zurückzugreifen, die sich für typische Einsatzszenarien adaptieren lassen. Ein Blueprint ist damit mehr als ein Codebeispiel: Er beschreibt Rollen, Kommunikationswege, Zustandsmodell, Tool-Integration, Sicherheitsgrenzen und Kontrollmechanismen als zusammenhängendes Systemdesign.

Im Kern beantworten Implementierungs-Blueprints fünf Fragen: Welche Agentenrollen werden benötigt? Wie interagieren diese Rollen miteinander? Welche Zustände und Daten müssen persistent gehalten werden? Welche externen Tools und APIs werden eingebunden? Und wie werden Sicherheit, Fehlertoleranz und menschliche Kontrolle umgesetzt? Erst wenn diese Ebenen zusammen gedacht werden, entsteht ein produktionsfähiges agentisches System.

Ein grundlegender Standard-Blueprint ist die Manager-Worker-Crew. Sie eignet sich für klar strukturierte, mehrstufige Aufgaben und besteht typischerweise aus einem koordinierenden Orchestrator sowie mehreren spezialisierten Ausführungsagenten. Der Manager übernimmt Task Decomposition, Priorisierung, Routing und Ergebnissynthese. Die Worker bearbeiten definierte Teilaufgaben wie Recherche, Analyse, Generierung oder Validierung. Ergänzt wird dieses Muster meist durch einen Review-Agent, der Resultate prüft und bei Bedarf Korrekturschleifen anstößt. Technisch empfiehlt sich für dieses Muster ein zentrales State-Objekt,

strukturierte Nachrichtenformate und ein gemeinsamer Zwischenspeicher für Artefakte.

Ein zweiter relevanter Blueprint ist die Research-and-Synthesis-Pipeline. Dieses Muster eignet sich für Marktforschung, Competitive Intelligence, Due Diligence oder Wissensarbeit. Typische Rollen sind Source-Discovery-Agent, Extractor-Agent, Analyst-Agent, Synthese-Agent und Quality-Reviewer. Der Prozess beginnt mit der Quellensuche, gefolgt von Extraktion und Normalisierung der Daten. Anschließend werden Muster, Trends oder Widersprüche identifiziert und in einer verdichteten Form zusammengeführt. Entscheidend ist hier ein sauberes Quellen- und Evidenzmodell, damit Aussagen nachvollziehbar bleiben und Halluzinationen minimiert werden.

Für operative Service-Prozesse eignet sich der Blueprint eines Support- und Eskalationssystems. Hier steht weniger die lineare Produktion eines Ergebnisses im Vordergrund als die fallbezogene Steuerung eines Workflows. Ein Intake-Agent klassifiziert Anfragen, ein Policy-Agent prüft Regeln und Berechtigungen, Fachagenten bearbeiten Spezialfälle, und ein Escalation-Agent entscheidet über Übergaben an menschliche Mitarbeiter. Dieses Muster benötigt ein robustes Ticket-State-Management, Zugriff auf Wissensdatenbanken sowie strikte Guardrails, damit Agenten nur zulässige Aktionen durchführen.

Ein besonders wirkungsvolles Muster im Entwicklungsumfeld ist die Engineering-Crew. Sie verbindet einen Planner-Agent mit Coding-, Testing-, Review- und Documentation-Agents. Der Planner übersetzt Anforderungen in Arbeitspakete, der Coding-Agent erstellt Implementierungen, der

Testing-Agent generiert Testfälle und prüft Laufverhalten, der Reviewer bewertet Qualität, Sicherheit und Stil, und der Documentation-Agent erzeugt technische Dokumentation. Dieses Muster funktioniert nur dann zuverlässig, wenn Tool-Use eng integriert ist: Repository-Zugriff, Test-Runner, Linter, statische Analyse und idealerweise eine isolierte Ausführungsumgebung.

Neben Rollenmodellen ist das State Design der kritische Kern jedes Blueprints. Jeder produktive Blueprint sollte mindestens zwischen drei Zustandsebenen unterscheiden: temporärer Arbeitskontext eines Agenten, geteilter Prozesszustand des Gesamtsystems und persistente Artefakte wie Ergebnisse, Logs oder Quellen. Ohne diese Trennung entstehen schnell Kontextverluste, unnötige Redundanz oder intransparente Entscheidungen. Empfehlenswert ist ein strukturiertes State-Schema mit Task-ID, Status, Verantwortlichem, Eingaben, Zwischenergebnissen, Risiken, Entscheidungen und Versionen.

Ebenso zentral ist die Kommunikationsarchitektur. Blueprints sollten festlegen, ob Agenten direkt per Messaging kommunizieren, über ein Blackboard-System interagieren oder einem graphbasierten Kontrollfluss folgen. Für einfache Systeme genügt häufig eine sequenzielle Übergabe mit strukturierten Outputs. Komplexere Umgebungen profitieren von Event-basierten Zustandsübergängen, insbesondere wenn mehrere Agenten parallel arbeiten oder sich Ergebnisse gegenseitig beeinflussen.

Kein Blueprint ist vollständig ohne Sicherheits- und Kontrollschicht. Dazu gehören Tool-Berechtigungen nach dem Least-Privilege-Prinzip, Laufzeitgrenzen, Retry-Limits, Abbruchkriterien für Schleifen, Logging, Audit-

Trails und klar definierte Human-in-the-Loop-Punkte. Besonders bei schreibenden oder transaktionalen Aktionen sollte jeder Blueprint zwischen lesenden, vorschlagenden und ausführenden Agenten unterscheiden. Was analysiert werden darf, sollte nicht automatisch auch ausgeführt werden dürfen.

Use Case	Empfohlene Topologie	Kritische Designpunkte
Content-Erstellung	Manager-Worker + Review	Klare Rollen, konsistente Prompt-Strukturen, iterative Qualitätssicherung
Marktforschung	Pipeline (Research → Analyse → Synthese)	Quellvalidierung, Evidenztracking, Kontextaggregation
Kundensupport	Hierarchisch + Eskalationslogik	Klassifikation, State Management, Human-in-the-Loop, Zugriffskontrolle
Software-Engineering	Manager-Worker + Tool-Integration	Code-Ausführung, Testing-Loops, Versionierung, Sandbox
Datenanalyse	Pipeline + Shared Memory	Datenkonsistenz, Normalisierung, Performance bei großen Datenmengen
Entscheidungsunterstützung	Multi-Agent Debate / Voting	Konsensmechanismen, Bias-Kontrolle, Ergebnisbewertung
Workflow-Automatisierung	Graphbasiert (z. B. Lang-Graph)	State-Transitions, Fehlerhandling, Wiederaufnahmefähigkeit
Research & Innovation	Peer-to-Peer + emergente Koordination	Exploration vs. Kontrolle, Iterationsgrenzen, Ergebnisvalidierung
Compliance & Audit	Hierarchisch + Policy-Engine	Nachvollziehbarkeit, Logging, Regelvalidierung, Zugriffsbeschränkungen
Vertrieb & Lead-Qualifizierung	Pipeline + Feedback-Loops	Datenqualität, Priorisierung, kontinuierliches Lernen

Für die praktische Umsetzung empfiehlt sich eine dreistufige Einführungslogik. In Phase eins wird ein Minimal-Blueprint als Prototyp gebaut, um Rollen, Prompts und Tool-Flows zu validieren. In Phase zwei folgt die Härtung: State Management, Monitoring, Fehlerbehandlung und Rechtekontrolle werden ergänzt. In Phase drei wird skaliert, indem zusätzliche Rollen, bessere Speichersysteme, Observability und Performanzoptimierung eingebaut werden. Viele Projekte scheitern, weil sie zu früh komplexe Agentenarchitekturen entwerfen, bevor ein schlanker Kern stabil funktioniert.

Der eigentliche Wert von Implementierungs-Blueprints liegt damit nicht in ihrer Wiederverwendbarkeit allein, sondern in ihrer Entscheidungsdisziplin. Sie zwingen dazu, Rollen, Datenflüsse, Zustände und Verantwortlichkeiten explizit zu machen. Genau darin liegt der Unterschied zwischen einer Demo und einem belastbaren Multi-Agenten-System. Voranstehende Matrix dient als schnelle Orientierung, um typische Anwendungsfälle mit passenden Architekturen und zentralen Designentscheidungen zu verknüpfen. Sie ersetzt keine detaillierte Systemplanung, bietet jedoch eine strukturierte Ausgangsbasis für die Auswahl geeigneter Blueprints.

Mehr zum Thema Multi-Agent-Systeme

Der vollständige Leitfaden „Synergie der Intelligenz“

 [Jetzt bei Amazon bestellen](#)