

Agentic AI auf einer Seite.

Was ein Agent ist, wann ihr einen braucht und was ihn im Zaum hält. Sechs Blöcke zum Speichern und Weitergeben, jeder mit dem Weg ins ausführliche Glossar.

01 Was ein Agent ist

03 Die Bausteine

05 Autonomie in Stufen

02 Skript, RPA oder Agent?

04 Der Governance-Ring

06 EU AI Act: die Einordnung

01 / 06

Was ein Agent ist

- Er bekommt ein **Ziel** statt einer Schrittfolge.
- Er arbeitet in einer Schleife, bis das Ziel erreicht ist oder eine Regel ihn stoppt.
- Chatbot antwortet · Copilot schlägt vor · **Agent führt aus.**
- Der Test: Wer legt die Schritte fest? Wählt das System Weg und Werkzeuge selbst, ist es ein Agent – was es ohne Freigabe ausführen darf, regeln die Autonomiestufen.



Im Glossar: KI-Agent · Agent vs. Chatbot vs. Copilot

02 / 06

Skript, RPA oder Agent?

Stabile Schnittstelle, feste Regeln

Skript oder API

Keine Schnittstelle, immer gleiche Masken, seltene Ausnahmen

RPA

Viele Varianten, unstrukturierte Dokumente, Urteil im Einzelfall

Agent

Kein Agent, solange es keine echten Fälle zum Testen gibt oder niemand für die Ausnahmen zuständig ist. In der Praxis verbindet ein Prozess oft alle drei – und Menschen.

03 / 06

Die Bausteine

MODELL

Versteht Sprache, garantiert aber nichts.

KONTEXT & RAG

Eure Dokumente als prüfbare Belege statt Raten.

TOOLS & MCP

Die Hände des Agenten. Eng geschnitten die wichtigste Kontrolle – Bildschirm- oder Code-Werkzeuge öffnen alles, was sein Konto darf.

EVALS

Qualität als Messwert auf echten Fällen – Ergebnis und Weg, mehrfach getestet.

04 / 06

Der Governance-Ring

IDENTITÄT

Eigene Agenten-Identität mit verantwortlicher Person, kurzlebige Tokens, nie ein geteilter Service-Account.

AUDIT

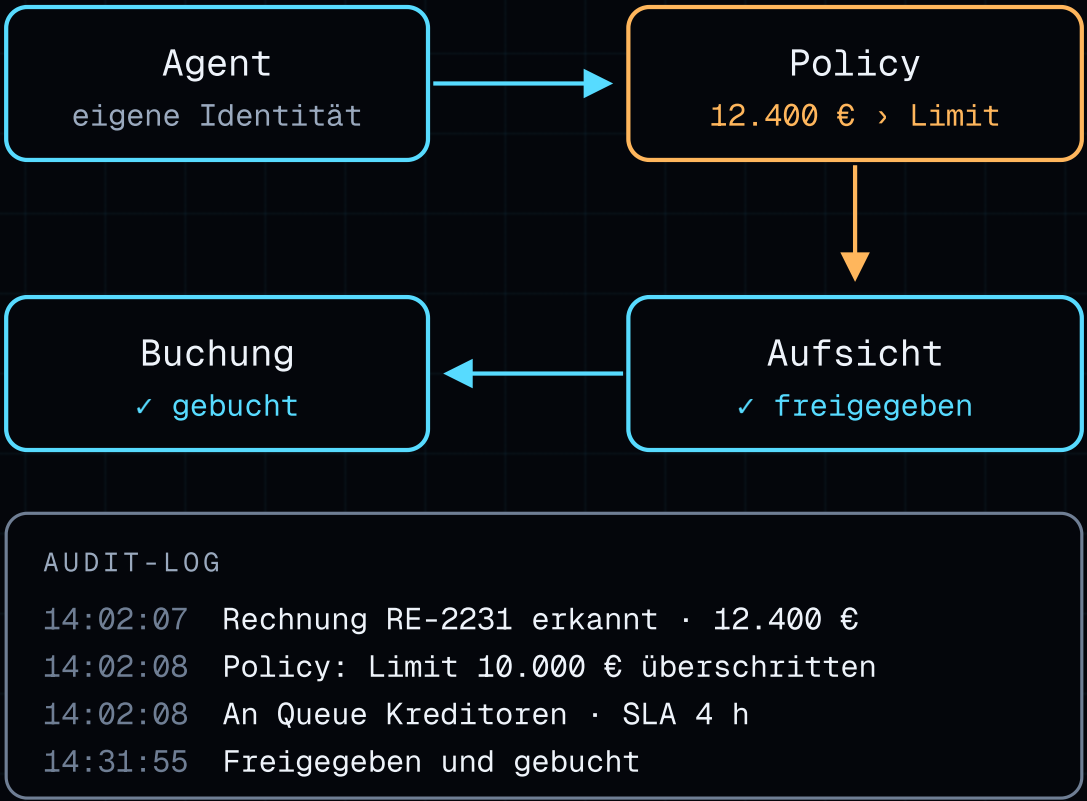
Jeder Schritt manipulationssicher protokolliert: Eingaben, Werkzeugaufrufe, Policy-Entscheid, Begründung.

POLICY

Grenzen stehen im Code und in der Tool-Schicht, nicht im Prompt – jede Mail, jedes PDF kann versteckte Anweisungen enthalten.

AUFSICHT

Ausnahmen landen in einer Queue mit SLA; der Stopp-Schalter greift außerhalb des Modells.



Im Glossar: Agenten-Identität · Guardrails & Policy · Audit-Trail · Human-in-the-Loop · Prompt Injection

05 / 06

Autonomie in Stufen

L0 Manuelle Arbeit

L1 RPA: feste Skripte

L2 ML im Prozess

L3 Agent schlägt vor, ein Mensch bestätigt

L4 Agent handelt in Grenzen, der Mensch greift bei Ausnahmen ein

Ein Reifemodell je Fallgruppe, keine Norm: Eine Stufe höher geht es erst, wenn Evals, Audit-Trail und Ausnahmequote es rechtfertigen.

Im Glossar: Autonomiestufen L0-L4

06 / 06

EU AI Act: die Einordnung



Die meisten Backoffice-Agenten liegen in den unteren beiden Stufen. Sie steigen auf, sobald ihre Entscheidungen Menschen direkt betreffen: Bewerbungen vorsortieren ist Hochrisiko.

seit 2. Aug. 2026	Transparenzpflichten nach Art. 50
2. Dez. 2026	Frist für maschinenlesbare Markierung bei Bestandssystemen
2. Dez. 2027	Hochrisiko nach Anhang III
2. Aug. 2028	Hochrisiko nach Anhang I

Stand September 2026, nach dem Digital Omnibus (VO (EU) 2026/1744). Keine Rechtsberatung.

Im Glossar: Risikoklassen · Der Zeitplan

Alle Begriffe ausführlich im Glossar.

[botforce-technology.com
/glossar](https://botforce-technology.com/glossar)

Vier Kapitel mit Beispielen aus der Praxis:
Automatisierung, KI, Agentic AI und EU AI Act.

Stand September 2026 · EU-AI-Act-Angaben
nach dem Digital Omnibus, keine
Rechtsberatung.