

Erste Schritte zur Nutzung des JURA KI Assistenten

(Stand 10.02.2026)

Inhalt

Was ist der JURA KI Assistent?	2
Wie starte ich den JURA KI Assistenten?.....	2
So nutzen Sie den JURA KI Assistenten in der Praxis.....	4
Prompt festlegen.....	4
Zu anonymisierenden Text einfügen.....	4
Externes KI-Sprachmodell auswählen	4
Google Gemini Modelle.....	5
Mistral AI-Modelle.....	6
OpenAI GPT-Modelle.....	7
Modelle von Perplexity	10
Hinweise zur Websuche	11
Anonymisieren und bearbeiten.....	11
Antwort von JURA KI	11
Was ist der JURA KI Publisher.....	12
So stellen Sie Entscheidungen bereit	12
Automatikmodus (empfohlen).....	12
Manuelles Hochladen einzelner Entscheidungen	12
Einstellungen im Überblick.....	12
Allgemein.....	12
JURA KI Assistent	13
Hilfen im Überblick.....	13

Was ist der JURA KI Assistent?

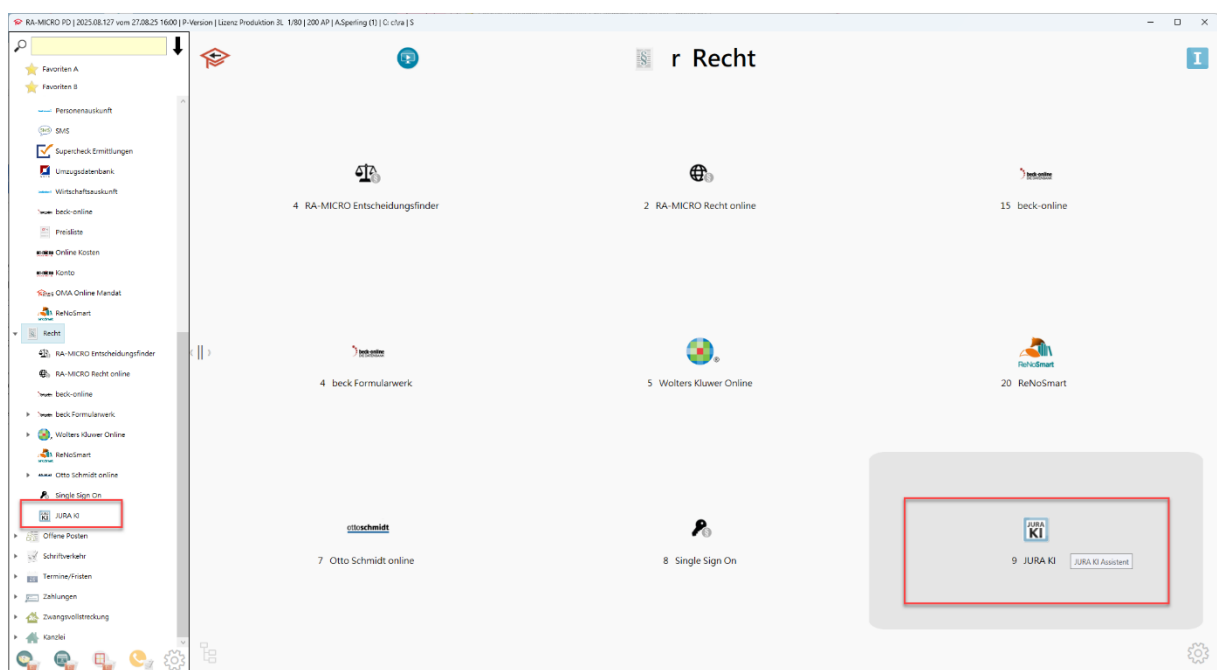
Der JURA KI Assistent unterstützt Sie im juristischen Arbeitsalltag mit leistungsstarken KI-Funktionen – speziell entwickelt für Anwaltskanzleien. Er ermöglicht u. a.:

- Datenschutzkonforme Anonymisierung von Texten mittels mitgelieferter interner KI, die lokal auf dem Kundengerät erfolgt
- Juristische Textverarbeitung (z. B. juristische Begründungen, Zusammenfassungen)
- Auswahl verschiedener externer KI-Modelle (OpenAI)
- Verlinkung und Verifizierung zitierter Vorschriften und Rechtsprechung

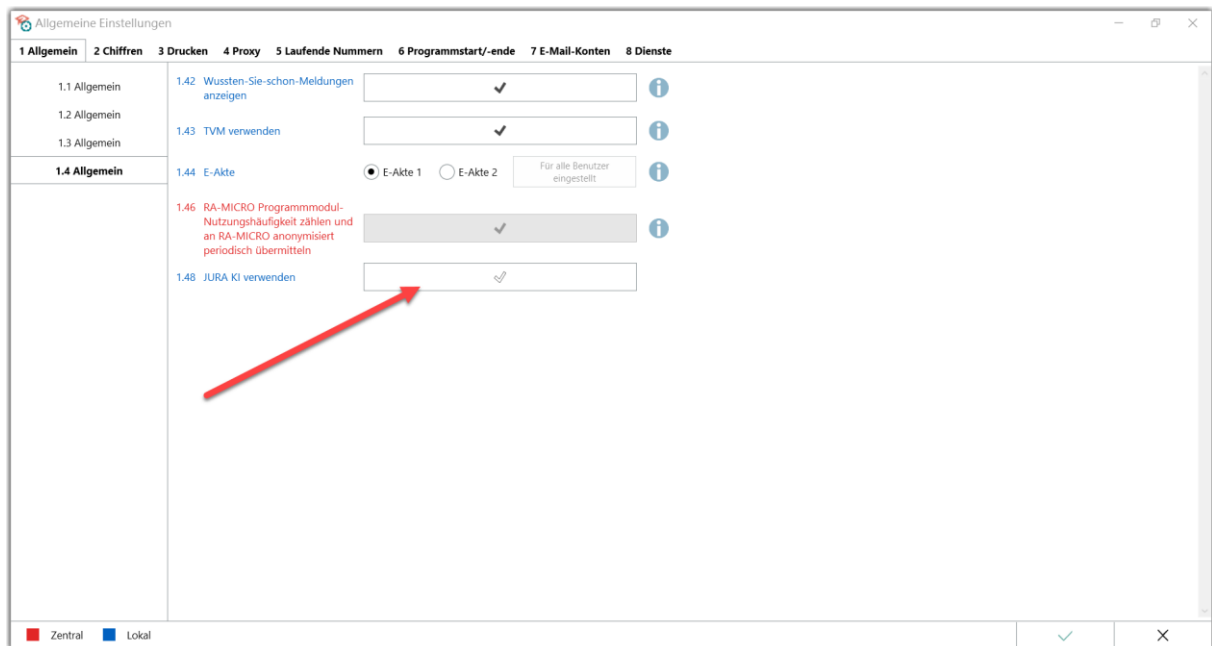
Wie starte ich den JURA KI Assistenten?

Als RA-MICRO Nutzer können Sie den JURA KI Assistenten direkt aus RA-MICRO heraus starten:

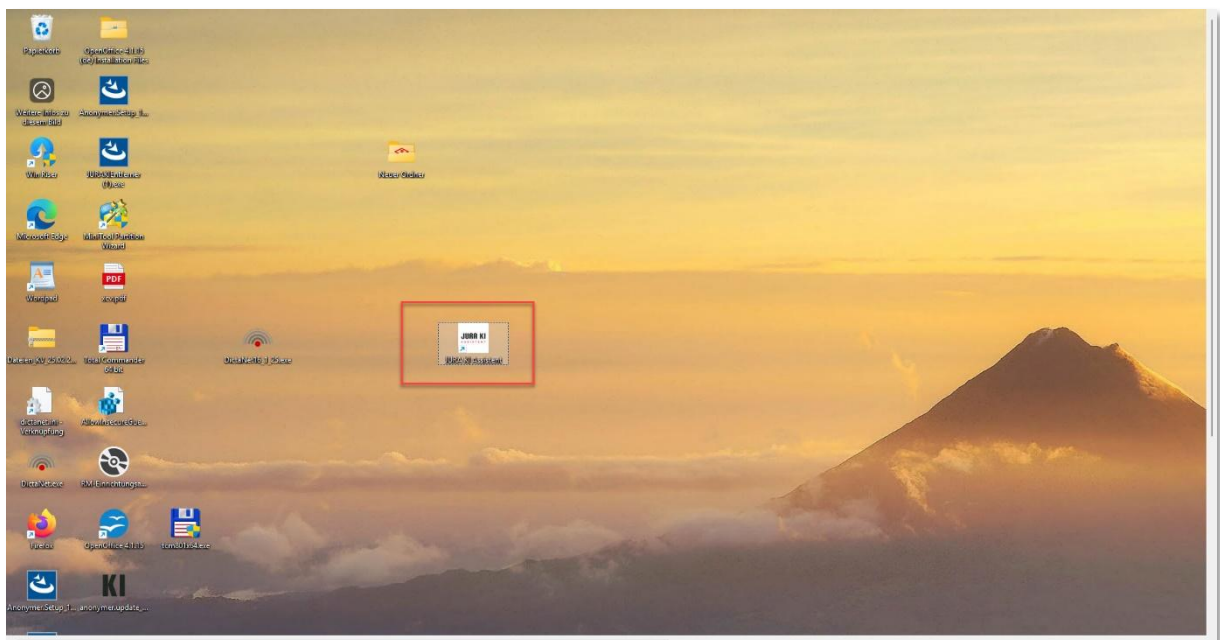
- Navigieren Sie zum **PD Recht** und **klicken** auf das **Icon des JURA KI Assistenten**



- Sollten Sie das Icon nicht sehen, aktivieren Sie zunächst den JURA KI Assistenten über **RA-MICRO / Allgemeine Einstellungen**



Haben Sie die Stand-Alone-Version des JURA KI Assistenten auf Ihrem Gerät installiert, starten Sie diesen durch **Doppelklick auf das Desktop-Icon**:



Eine ausführliche Anleitung zur Installation und Ersteinrichtung des JURA KI Assistenten finden Sie unter <https://ki-lab.dictanet.com/downloads/Anleitung-Installation-Ersteinrichtung-ausfuehrlich.pdf>.

So nutzen Sie den JURA KI Assistenten in der Praxis

Prompt festlegen

- Wählen Sie im oberen linken Bereich einen passenden Prompt aus den **Allgemeinen Promptvorlagen** oder erstellen Sie Ihren **eigenen Prompt**. Sie können dort auch **eigene Promptvorlagen** erstellen und verwalten.
- Der Prompt ist die Arbeitsanweisung an JURA KI und bestimmt, wie die KI Ihre Aufgabe interpretiert – strukturiert, zusammenfassend und/oder mit juristischer Argumentation.
- Rechts neben der Promptauswahl können Sie erforderlichenfalls zusätzlich eine **Sprache** auswählen, in welcher die Antwort zusätzlich automatisch übersetzt wird.

Tipp – Bessere Antworten durch Prompt Optimierung

Unsere mitgelieferten Vorlagen sind als Orientierung gedacht. Wir empfehlen, diese nach Bedarf zu verfeinern und mit unterschiedlichen Modellen zu testen, um die für Ihren Anwendungsfall besten Ergebnisse zu erzielen.

Zu anonymisierenden Text einfügen

- Fügen Sie im Feld **Zu anonymisierender Text** Ihren Sachverhalt oder Ihre zu verarbeitenden Dokumente ein. Fügen Sie bei Bedarf Texte via Copy & Paste-Funktion ein oder laden Sie ganze Schriftstücke hoch, die verarbeitet werden sollen.
- Unter **Jur. Begründung** finden Sie bereits einen beispielhaften Text, der als Orientierung dienen kann.

Externes KI-Sprachmodell auswählen

- Im unteren linken Bereich finden Sie die Modellauswahl. Wählen Sie das gewünschte externe Sprachmodell, welches die Anfrage bearbeiten soll. Folgende Modelle stehen aktuell zur Verfügung:

Google Gemini Modelle

Gemini 3.0 Pro

Das Google-Modell **Gemini 3.0 Pro** ist Teil der neuesten Generation der Gemini Architektur und bietet modernste Schlussfolgerungsfähigkeit, multimodale Intelligenz und einen extrem großen Kontext, basierend auf den neuesten Fortschritten von Google AI/DeepMind. Es wurde konzipiert, um anspruchsvolle Aufgaben mit tiefer Analyse und umfassendem Wissen zu bewältigen.

Vorteile:

- Verbesserte Leistung bei komplexem Reasoning, tiefgehenden Analysen und strukturierten Argumentationen – ideal für juristische Fachfragen
- Großes Kontextverständnis – verarbeitet sehr lange Texte kohärent
- Leistungsstarke multimodale Fähigkeiten für umfassende Auswertungen

Einsatzempfehlung:

Empfohlen für tiefgehende juristische Analysen, kritische Schriftsätze oder die Auswertung großer Dokumente, wenn höchste Präzision, tiefgründige Argumentation und fundierte Ergebnisse gefragt sind.

Gemini 3.0 Flash

Das Google-Modell **Gemini 3.0 Flash** ist eine auf Geschwindigkeit, Effizienz und Wirtschaftlichkeit optimierte Variante der Gemini 3 Generation. Es kombiniert die Intelligenz und Reasoning Fähigkeiten des Pro Modells mit deutlich geringerer Latenz und besseren Kosten /Leistungswerten.

Vorteile:

- Schnelle Verarbeitung und sehr kurze Antwortzeiten – ideal für effiziente Nutzung im juristischen Alltag
- Gutes Gleichgewicht aus Intelligenz, Reasoning und Ressourceneffizienz – für Standard und mittelkomplexe Aufgaben
- Basierend auf derselben leistungsfähigen Reasoning Plattform wie Gemini 3 Pro, mit Steuerung des Denkniveaus zur Balance zwischen Aufwand und Ergebnisqualität

Einsatzempfehlung:

Optimal für juristische Standard Analysen, schnelle Sachverhaltsauswertungen, Routine Recherchen und effiziente Schriftsatzkorrekturen, insbesondere dort, wo Geschwindigkeit und Kostenbewusstsein im Vordergrund stehen.

Gemini 2.5 Pro

Das Google-Modell **Gemini 2.5** basiert auf der neuesten Generation der Gemini-Architektur und ist auf komplexes Denken, Analyse und Argumentationsführung ausgerichtet. Es eignet sich hervorragend für tiefgreifende juristische Auswertungen mit umfangreichem Kontext.

Vorteile:

- Sehr hohe Textqualität bei anspruchsvollen Aufgaben
- Geeignet für längere Eingaben mit mehrschichtiger Argumentation
- Starke Leistung bei Gutachten, Urteilsanalysen und logischen Prüfungsketten

Einsatzempfehlung:

Ideal für juristische Fachanalysen, komplexe Schriftsatzprüfungen und strukturierte Begründungen mit hoher Detailtiefe – insbesondere bei anspruchsvollem fachlichem Niveau.

Gemini 2.5 Flash

Das Google-Modell **Gemini 2.5 Flash** ist eine performante, besonders schnelle Variante der Gemini-Architektur. Es bietet sehr gute Textqualität bei deutlich reduzierter Antwortzeit – ideal für Aufgaben mit mittlerer Komplexität.

Vorteile:

- Schnelle und zuverlässige Textverarbeitung für typische Arbeitsabläufe
- Sehr gutes Gleichgewicht aus Qualität, Geschwindigkeit und Kosten
- Solide Argumentationsstruktur bei juristischen Standardfällen

Einsatzempfehlung:

Optimal für Kanzleialltag, Entwurfshilfen, rechtliche Kurzanalysen und einfache Schriftsatzkorrekturen mit klarem Fokus auf Effizienz.

Mistral AI-Modelle**Mistral Small 3.2**

Das AI-Modell **Mistral Small 3.2** ist eine kompakte, optimierte Modellvariante für juristische Aufgaben mit mittlerer Komplexität. Es kombiniert hohe Antwortqualität mit robuster Instruktionsverarbeitung – bei gleichzeitig schlankem Ressourcenverbrauch.

Vorteile:

- Zuverlässige Textverarbeitung für typische juristische Standardprozesse
- Gute Balance zwischen Argumentationsqualität, Modellgröße und Rechenaufwand
- Flexibel einsetzbar für verschiedene Textarten und -längen

Einsatzempfehlung:

Geeignet für allgemeine juristische Textverarbeitung, Entwurfshilfen und einfache Analysen – besonders bei regelmäßigem Einsatz mit Fokus auf Wirtschaftlichkeit.

Mistral Medium 3.1

Das Modell **Mistral Medium 3.1** ist die leistungsstärkste Variante der Medium-Reihe von Mistral AI und bietet sehr gute Textqualität für juristische Aufgaben mit hohem Anspruchsniveau. Es ist auf lange Kontexte und strukturierte Argumentationen ausgelegt.

Vorteile:

- Hervorragende Textqualität bei anspruchsvollen Aufgaben
- Geeignet für längere Eingaben mit komplexen Anforderungen
- Starke Leistung bei juristischen Analysen und Argumentationen

Einsatzempfehlung:

Optimal für detaillierte juristische Analysen, komplexe Schriftsatzprüfungen und strukturierte Begründungen – insbesondere bei höherem fachlichen Anspruch.

Hinweis zur Nutzung:

Dieses Modell ist mit dem **kostenfreien** Tarif **Experiment nicht verfügbar**.

Mistral Medium 3

Das Modell **Mistral Medium 3** ist eine bewährte Variante der Medium-Reihe mit starker Leistung bei juristischen Anwendungen. Es unterstützt umfangreiche Textverarbeitung und eignet sich gut für anspruchsvolle, aber standardisierte Aufgaben.

Vorteile:

- Solide Argumentationsstruktur bei Gutachten und längeren Schriftsätzen
- Zuverlässige Verarbeitung juristischer Texte mit mittlerer bis hoher Komplexität
- Kosteneffizienter Einsatz bei regelmäßigem Bedarf an hochwertiger Textausgabe

Einsatzempfehlung:

Empfohlen für juristische Auswertungen, komplexe Schriftsatzentwürfe oder Mandatsvorbereitungen mit größerem Umfang – insbesondere bei regelmäßigem Einsatz und Fokus auf Qualität bei stabilem Ressourcenrahmen.

Hinweis zur Nutzung:

Dieses Modell ist mit dem **kostenfreien** Tarif **Experiment nicht verfügbar**.

OpenAI GPT-Modelle

GPT-5.2

Das OpenAI-Modell **GPT-5.2** ist das aktuell leistungsstärkste Modell der GPT-5-Serie. Es liefert deutliche Verbesserungen der Allgemeinintelligenz, dem Lang-Kontext-Verständnis, ein mehrstufiges Reasoning sowie reduzierte Fehlerquoten.

Vorteile:

- Start erweitertes Verständnis langer, komplexer Sachverhalte – ideal bei umfangreichen Dokumenten
- Verbessertes mehrstufiges Reasoning und strukturierte Argumentationsausgabe bei juristischen Analysen
- Höhere Genauigkeit und robustere Befolgung juristischer Anweisungen – weniger Halluzinationen / Fehlantworten

Einsatzempfehlung:

Empfohlen für tiefgehende juristische Analysen, komplexe Dokumentenauswertung, umfangreiche Gutachten oder kombinierte Aufgaben – insbesondere dort, wo maximale Präzision, kohärente Lang-Kontexte und stabile Argumentationsketten gefordert werden.

GPT-5.1

Das OpenAI-Modell **GPT-5.1** ist ein gezieltes Update der GPT-Reihe, welches die Intelligenz, Gesprächsqualität und adaptive Reasoning Fähigkeiten gegenüber früheren GPT-Modellen weiter verbessert.

Vorteile:

- Verbessertes adaptives Reasoning – das Modell entscheidet, wie „tief“ es über eine Anfrage nachdenken muss
- Höhere Genauigkeit bei juristischen Texten und strukturierter Argumentation
- Natürlicher, dialogorientierter Kommunikationsstil (klarer, leichter steuerbar)

Einsatzempfehlung:

Optimal für juristische Standard Analysen, strukturierte Dokumentenverarbeitung und Mandantenkommunikation, wenn es auf ein gutes Gleichgewicht aus Antwortqualität, Verständlichkeit und Effizienz ankommt – z. B. bei Schriftsatzprüfungen, Fallzusammenfassungen und Argumentationsaufbau im Kanzleialltag.

GPT-4.1

Das OpenAI-Modell **GPT-4.1** basiert auf der verbesserten GPT-4-Architektur und bietet solide Sprachverarbeitung bei sehr guter Textqualität. Es eignet sich für vielseitige juristische Anwendungsfälle mit mittlerem bis hohem Anspruchsniveau.

Vorteile:

- Zuverlässige Textverarbeitung mit guter Balance aus Qualität, Tempo und Wirtschaftlichkeit
- Gut geeignet für juristische Textprüfungen, Analysen und Entwürfe
- Unterstützt längere Eingaben mit fundierter Argumentation

Einsatzempfehlung:

Ermöglicht Schriftsatzprüfungen, rechtliche Argumentationen und Mandantenkommunikation auf gutem fachlichem Niveau – auch bei regelmäßigem Einsatz.

GPT-4.1 mini

Das OpenAI-Modell **GPT-4.1 mini** ist eine kompaktere Variante mit kürzeren Reaktionszeiten und reduziertem Rechenaufwand. Es eignet sich besonders für standardisierte Aufgaben mit klarem Fokus auf Effizienz.

Vorteile:

- Rasche Verarbeitung – auch bei mehreren gleichzeitigen Abfragen
- Wirtschaftlich im Betrieb – optimal für skalierte Nutzung
- Gute Textqualität für klassische Kanzlei- oder Verwaltungstätigkeiten

Einsatzempfehlung:

Empfehlenswert für standardisierte Abläufe im juristischen Arbeitsalltag – wie etwa Formulierungshilfen, Vermerke, allgemeine Schriftstücke oder einfache Prüfaufträge.

GPT-4.1 nano

Das OpenAI-Modell **GPT-4.1 nano** ist auf maximale Geschwindigkeit und Effizienz ausgelegt und bietet eine stabile Textqualität für einfache Aufgaben mit geringem Kontextbedarf.

Vorteile:

- Sehr schnelle Reaktionszeiten – ideal für einfache Prüfanfragen, Kurzantworten oder Textbausteine
- Geringe Systembelastung – ermöglicht flüssige Nutzung auch bei hoher Abfragezahl
- Stabile Textqualität bei reduzierter Tiefe – ausgelegt auf kurze, direkte Ausgaben

Einsatzempfehlung:

Optimal für häufig wiederkehrende, klar umrissene Aufgaben – z. B. bei einfachen Prüfanfragen, der Generierung von Textbausteinen oder der schnellen Erstbeantwortung juristischer Rückfragen.

GPT-4o

Das OpenAI-Modell **GPT-4o** ist das aktuelle Multimodal-Flaggschiff von OpenAI und bietet höchste Antwortqualität bei gleichzeitig sehr guter Performance.

Vorteile:

- Schnell und kosteneffizient: Deutlich schneller als frühere GPT-4-Varianten bei geringerem Ressourcenverbrauch

- Zuverlässige Textqualität: Bei juristischen Fragestellungen liefert GPT-4o strukturierte, nachvollziehbare Ergebnisse
- Große Kontexte: Unterstützt Eingaben und Ausgaben mit bis zu 128.000 Tokens – ideal für lange Dokumente, Schriftsätze oder Gutachten

Einsatzempfehlung:

Für anspruchsvolle Standardaufgaben ebenso geeignet wie für umfangreiche Texteingaben. Besonders hilfreich bei juristischen Analysen, intelligenten Zusammenfassungen oder der Verarbeitung großer Textmengen mit kontextabhängigen Anforderungen.

GPT-4o mini

Das Modell GPT-4o mini ist eine besonders effiziente Variante des GPT-4o-Modells und bietet hohe Geschwindigkeit bei geringem Ressourcenverbrauch.

Vorteile:

- Schnell und effizient: Ideal für einfache, wiederkehrende Aufgaben mit geringem Ressourcenbedarf
- Kosteneffizient: Spart Tokens im Vergleich zum vollwertigen GPT-4o-Modell
- Solide Qualität: Liefert verlässliche Ergebnisse bei typischen Kanzleifällen, Textbausteinen oder einfachen Analysen

Einsatzempfehlung:

Ideal bei hohem Anfragevolumen, für wiederholte Interaktionen oder standardisierte Abläufe – etwa für Formulierungshilfen, Vorlagen, einfache Prüfanfragen oder interne Kommunikation mit reduziertem Kontextumfang.

Modelle von Perplexity

Perplexity Sonar

Das Modell **Sonar** von Perplexity ist auf schnelle, webbasierte Rechercheprozesse ausgerichtet. Es liefert kompakte juristische Kurzantworten mit aktuellen Quellenbezügen aus dem Internet – inklusive strukturierter Ausgabe.

Vorteile:

- Sehr schnelle Reaktionszeit für juristische Kurzrecherchen
- Integrierte Quellenangaben für nachvollziehbare Ergebnisse
- Zugriff auf tagesaktuelle Informationen bei rechtlichem Webbezug

Einsatzempfehlung:

Besonders geeignet für juristische Recherchen mit Aktualitätsbezug – z. B. Gesetzesänderungen, neue Urteile oder Medienberichte. Ideal bei kurzen Fragestellungen mit Fokus auf schnellen Zugriff und Faktenprüfung.

Hinweise zur Websuche

Die Websuche ermöglicht dem Modell, aktuelle Informationen aus dem Internet abzurufen und – je nach Anbieter/Modell – Quellen zu nennen. Wenn Websuche für das ausgewählte Modell verfügbar ist, kann Sie neben der Modellauswahl aktiviert oder deaktiviert werden.

Manche Modelle haben anbieterbedingt Websuche immer aktiviert oder unterstützen sie nicht.

Websuche kann Antworten aktueller machen, ist aber oft langsamer und kann bei hoher Auslastung oder umfangreichen Recherchen zu Einschränkungen/Fehlern führen (z. B. Rate-Limits/429). Für stabile Aufgaben ohne Aktualitätsbezug empfehlen wir, die Websuche auszuschalten (schneller/stabiler).

Anonymisieren und bearbeiten

Die Anonymisierung erfolgt datenschutzkonform durch ein **im JURA KI Assistenten integriertes, lokales KI-Modell** automatisiert direkt auf Ihrem Gerät, ohne dass Daten an externe Server übermittelt oder dauerhaft gespeichert werden.

Die Aufgabe wird sortiert an das externe Modell via API-Zugriff übergeben, sodass dieses genau weiß, wie es reagieren soll – klar, strukturiert und datenschutzkonform.

Antwort von JURA KI

- Die Antwort des externen Modells wird wiederum lokal auf Ihrem Gerät durch den JURA KI Assistenten aufbereitet und deanonymisiert im Feld **Antwort von JURA KI** zur Verfügung gestellt.
- Verifizierte Fundstellen werden übersichtlich als **Liste der enthaltenen Links** ergänzt, sodass Sie schnell auf die relevanten Quellen zugreifen können.
- Bei unklaren Antworten kann ein **Ergänzender Prompt** hilfreich sein: Über diesen können Sie gezielte Rückfragen oder zusätzliche Anweisungen im Kontext stellen und die Antwort so überarbeiten lassen.
- Wählen Sie **Herunterladen**, wird der Antworttext automatisch in der Zwischenablage zur Weiterverarbeitung abgelegt. Zusätzlich wird für jede ausgeführte Aufgabe eine **Job-Datei** (wahlweise *.docx, *.pdf, *.txt) generiert und in Ihrem Downloadverzeichnis abgelegt.
- Über den Toggle **Nur deanonymisierte Antwort speichern** stellen Sie ein, ob nur der reine, aufbereitete Antworttext oder die gesamte Aufgabe inkl. Prompt, anonymisiertem Sachverhalt und anonymisierter Antwort in der Job-Datei abgebildet wird.

Was ist der JURA KI Publisher

Der JURA KI Publisher erweitert den JURA KI Assistenten um eine innovative Funktion zur kollaborativen Stärkung der juristischen Datenbasis in Deutschland. Ziel ist es, nicht publizierte Gerichtsentscheidungen aus Anwaltskanzleien automatisiert zu anonymisieren und in einer zentralen Rechtsprechungsdatenbank zusammenzuführen. Diese bildet die Grundlage für qualitativ hochwertige KI-Antworten. Der JURA KI Publisher bietet:

- Einfache Beteiligung durch Auswahl eines lokalen Dokumentenordners mit Entscheidungen
- Lokale, datenschutzkonforme Anonymisierung durch interne KI
- Automatische Erkennung relevanter Urteile (z. B. mit "Im Namen des Volkes")
- Upload in die RA-MICRO KI Rechtsprechungsdatenbank
- Nutzung der Inhalte für sogenannte Vorschaltmodelle, die die Qualität juristischer KI-Antworten deutlich verbessern

So stellen Sie Entscheidungen bereit

Automatikmodus (empfohlen)

- Aktivieren Sie die Checkbox **Automatikmodus**.
- Geben Sie anschließend den Datenpfad ein, in dem die Dokumente der Kanzlei inkl. Gerichtsentscheidungen gespeichert sind.

Manuelles Hochladen einzelner Entscheidungen

- Fügen Sie den Entscheidungstext direkt aus der Zwischenablage ein oder laden Sie eine Datei hoch.
- Ergänzen Sie das **Aktenzeichen** und klicken Sie auf **Hochladen** – fertig.
- **Optional:** Tragen Sie weitere Angaben wie Rechtsebene, Einzelschriften, Schlagwörter oder persönliche Kommentare ein. Diese Informationen verbessern die Verarbeitung in der Datenbank und die Qualität der daraus entstehenden Vorschaltmodelle.

Einstellungen im Überblick

Allgemein

- Ihre RA-MICRO Online Benutzerdaten (Kundennummer, Benutzerkennung)
- Hinterlegung Ihres persönlichen API-Keys zur Verbindung mit OpenAI
- Optional: Einrichtung eines Proxyservers für die Netzwerkverbindung

JURA KI Assistent

- Auswahl der Anonymisierungsmethode (BIO-Tags oder Einzelbuchstaben)
- Definition des Umgangs mit Gerichtsentscheidungen (Rubrum entfernen oder anonymisieren)
- Festlegen, wie mit nicht verifizierbaren Zitaten in KI-Antworten umgegangen werden soll
- Festlegen, wie die Antwort von JURA KI zur Weiterverarbeitung ausgegeben werden soll
- Zusätzliche Einstellungen für Schriftart und -größe bei der Verwendung von Word-Dokumenten als Ausgabedatei

Hilfen im Überblick

- Zugriff auf Anleitungen und Konfigurationen
- Verlinkung zum JURA KI Wiki mit häufig gestellten Fragen (FAQ)
- Einblick in aktuelle Versionshinweise und Updates