

Anmerkung: Dieses Dokument ist mit Überschriften gegliedert. JAWS-Anwender können durch Drücken der Taste H von Überschrift zu Überschrift navigieren. Alternativ kannst du die JAWS-Taste + F6 drücken, um eine Übersicht aller Überschriften zu erhalten und direkt zum gewünschten Punkt zu springen.

Einleitung

Thaliruth Arda Intelligence

Das Programm bedienen: Handbuch zu allen Bereichen

Kunden-Dokumentation, Teil 3, vollständige Fassung, Stand 01. August 2026, passend zu Programmversion 0.8.228. Dieses Dokument wird stets bei Änderungen oder Neuerungen ergänzt.

Auf den ersten acht Seiten erkläre ich Details zum Lizenzmodell, dem Geschäftsmodell sowie einige persönliche Hintergründe zu meiner Person. Da ich von Grund auf ein sehr transparenter Mensch bin und großen Wert darauf lege, keine Geheimnisse zu haben, möchte ich, dass jeder meine

Absichten – auch im Hinblick auf die Preisgestaltung – nachvollziehen kann. Wer sich direkt mit dem Umfang des Programms befassen möchte, navigiert bitte zur Seite sechs unter der Überschrift:

1. Über dieses Handbuch.

Audio-Inhalte und Apple Podcasts

Die unten aufgeführten Abschnitte werden in Kürze durch Links zu Apple Podcasts ergänzt. Dort kannst du dir Folgen zu den jeweiligen Themen anhören, wodurch du direkt erleben kannst, wie die Nutzung mit JAWS funktioniert.

Anmerkung 1: Status der Links

Die Links folgen erst später, da die entsprechenden Podcast-Episoden momentan noch nicht aufgezeichnet wurden. Dieses Dokument dient primär dazu, dir die Funktionen und den Überblick über das Programm zu präsentieren.

Preisgestaltung und Lizenzmodell

Die endgültigen Preise stehen noch nicht fest. Thaliruth Arda Intelligence (TAI) wird eine Basic-Version bieten, welche die Standardfunktionen für Chat und Projekte umfasst. Zusätzliche Bereiche können optional als Erweiterung dazugebucht werden, oder du entscheidest dich direkt für die Premium-Variante, die dir den vollen Funktionsumfang bietet.

Das Lizenzsystem von TAI basiert auf einem lebenslangen Schlüssel, der auf bis zu drei Geräten aktiviert werden kann. Solltest du mehr Geräte

benötigen, kannst du dies beim Kundendienst anfragen. Im Kundenbereich kannst du deine registrierten Geräte jederzeit verwalten und beispielsweise alte Computer entfernen, wenn du ein neues Gerät in Betrieb nimmst. Wenn du die Basic-Version besitzt und später Erweiterungen wie Hörspiele, Podcasts oder die Hörspiel-Werkstatt buchst, wird dies direkt auf deiner bestehenden Lizenz aktiviert; du benötigst also keinen neuen Lizenzschlüssel.

Testphase und Kaufprozess

Vor dem Kauf steht dir eine 7-tägige Demo-Version mit einer Zeitlizenz zur Verfügung. (Diese Zeitspanne ist Aktuell nur ein Platzhalter, möglicherweise wird es auch eine Testphase von 14 bis 30 Tagen geben. Dies wird noch Entschieden.) Diese bietet dir vollen Zugriff auf alle Premium-Funktionen, sodass du das Programm in Ruhe testen kannst. Bitte beachte, dass die Nutzung nach Ablauf der 7 Tage nicht mehr möglich ist. Falls du dich für einen Kauf entscheidest, wird deine bestehende Lizenz einfach umgestellt. Ein erneutes Eingeben eines neuen Lizenzschlüssels ist nicht nötig; klicke dazu nach dem Programmstart im Hinweisfenster einfach auf „Lizenz erneut prüfen“.

Preisgestaltung und Datenschutz

Die Basic-Version wird zum Veröffentlichungstermin einen Einführungspreis unter 100 Euro haben, der später auf über 100 Euro steigen wird. Mir ist bewusst, dass dies für die Basic-Version eine Investition darstellt. Man sollte jedoch bedenken, dass es auf dem Markt derzeit kein vergleichbares barrierefreies Programm von einem blinden KI unterstützten Entwickler für blinde Menschen gibt. Zumindest was meine Recherchen ergaben, habe ich nichts vergleichbares finden können. Will es aber grundsätzlich nicht Ausschließen! Ein entscheidender Vorteil ist zudem der Datenschutz: TAI kann komplett lokal und ohne Internetzugang betrieben werden, sodass alle sensiblen Daten auf deinem eigenen Rechner bleiben.

Vergleich: Cloud-basierte Dienste vs. Lokale Lösung

Ein Beispiel verdeutlicht das Problem bei Online-Diensten: Wenn du in JAWS die Funktion „Sprechendes Bild“ nutzt, um einen Screenshot analysieren zu lassen, wird dieses Bild an einen Server in den USA gesendet und dort ausgewertet. Dabei besteht Unklarheit darüber, was mit der Aufnahme geschieht – etwa ob sie zum Training von KI-Modellen verwendet wird. Zudem ist bekannt, dass JAWS künftig verstärkt KI-Funktionen und verpflichtende Benutzerkonten integrieren soll, was die Datentransparenz einschränkt. Große Unternehmen profitieren massiv von den gesammelten Nutzerdaten, während die eigentlichen Datengeber leer ausgehen.

Im Gegensatz dazu nutzt du bei Thaliruth Arda Intelligence einfach die Tastenkombination `STRG + SHIFT + B`. Ein Screenshot wird erstellt und direkt auf deinem PC durch die KI ausgewertet. Alle Informationen bleiben lokal. Du entscheidest selbst über den Nutzen dieser Funktion. Selbst im Vergleich zu einem Abo wie ChatGPT Premium (ca. 23 Euro monatlich) hätte sich die Kostenersparnis bereits nach fünf Monaten amortisiert – und du hältst dauerhaft eine lokale, barrierefreie Lösung ohne laufende Kosten. Bei Problemen stehe ich dir zudem unterstützend zur Seite.

Und wer mich kennt, der weiß sehr gut, das ich persönlich stets viel Wert auf Transparenz lege, und das seit über 15 Jahren Community Arbeit.

Entwicklungskosten und Geschäftsmodell

Ich habe etwa ein halbes Jahr an diesem Projekt gearbeitet und die Umsetzung mithilfe von KI realisiert. Allein die Kosten für die KI belaufen

sich auf rund 550 Euro. Rechnet man Arbeitszeit und Stromkosten hinzu, wären die Preise, die ich verlangen müsste, kaum zumutbar.

Es handelt sich um einen Einmalkauf, kein Abonnement! Das einzige Abo-Element ist die Funktion für Updatepakete, wobei auch hier die Entscheidung bei dir liegt. Sicherheitsupdates sind jederzeit kostenlos verfügbar. Wenn du ein Updatepaket für drei Monate erwirbst, erhältst du in diesem Zeitraum alle neuen Funktionen und Verbesserungen. Nach Ablauf der drei Monate erhältst du weiterhin Sicherheitsupdates, aber keine neuen Features mehr. Das Programm bleibt jedoch gemäß deiner Lizenz uneingeschränkt nutzbar. Ich lege zudem Wert auf ein schlankes Programm und werde es nicht unnötig weiter mit Funktionen aufblähen.

Auch werde ich nicht wirklich viel Gewinn erwirtschaften, da laufende Betriebskosten wie Server- und Stromkosten sowie meine eigene Arbeitszeit anfallen. Allein durch den Verkauf einzelner Lizenzen werde ich in den ersten Monaten und Jahren keine schwarzen Zahlen schreiben.

Wenn beispielsweise in einem Monat eine Basic-Version bestellt wird, verzeichne ich nach Abzug der laufenden Kosten allein durch diesen einen Verkauf immer noch ein Minus von etwa 250 Euro. Das bedeutet, dass das Projekt privat finanziert wird. Zudem kann es vorkommen, dass ich in einem Monat ein Programm für 150 Euro verkaufe und danach sechs Monate lang gar nichts verkaufe. Selbstständigkeit bedeutet im Gegensatz zu einem festen Arbeitsplatz kein gesichertes Einkommen. Dies sollte jeder berücksichtigen, bevor er den Preis kritisiert, ohne die dahinterstehende Infrastruktur und die Betriebskosten zu kennen. Mir geht es nicht primär darum, Geld zu verdienen; in erster Linie möchte ich mit meiner Software anderen Sehbeeinträchtigten helfen. Wenn natürlich ein Gewinn übrig bleibt, ist das hervorragend, da dies mir zeigt, dass meine Arbeit geschätzt und akzeptiert wird. Die Preisgestaltung bei Software weist stets eine große Spannweite auf, je nach Entwicklungsaufwand, Entwicklungskosten und der benötigten Zeit: Es gibt Apps und Cp, Computer-Anwendungen für unter zehn Euro, aber ebenso solche, die hunderte oder sogar tausende Euro kosten können.

Mein Ziel ist es, mit meinen Programmen anderen Blinden zu helfen und ihnen Werkzeuge an die Hand zu geben, die es so noch nicht gibt. Diesen Mangel habe ich seit meiner Erblindung selbst schmerzlich erfahren müssen. Genau deshalb habe ich beschlossen, selbst aktiv zu werden und die Software zu entwickeln, die ich benötige, um auch unabhängiger von anderen Diensten zu werden. Ich bin überzeugt: Programme, die mir helfen, helfen auch anderen!

Darüber hinaus werde ich auf Anfrage individuelle Softwarelösungen anbieten. Wenn ein Kunde einen speziellen Bedarf äußert, erstelle ich einen Plan und setze diesen um, sodass der Kunde am Ende ein perfekt auf ihn abgestimmtes Programm erhält.

BlindSoft wird nicht nur Software produzieren, sondern auch barrierefreie Weblösungen realisieren, wie etwa eine Plattform für Podcast-Hosting oder schlichte Webauftritte. Und wie der Name BlindSoft bereits andeutet, besteht auch Raum für internationale Lösungen.

Denn ich entwickle keine barrierearme, sondern echte barrierefreie Software!

Ich bin kein gelernter Softwareentwickler, sondern Koch! Da ich jedoch bereits seit den 1990er-Jahren mit Computern arbeite und von 2000 bis 2010 selbstständig im IT-Bereich tätig war, verfüge ich über fundierte Kenntnisse in diesem Umfeld. Mein damaliges Gewerbe habe ich aufgegeben – unter anderem aufgrund des hohen Konkurrenzdrucks und vor allem, weil ich zu oft meinem Honorar hinterherlaufen musste.

Kunden erhielten monatliche Rechnungen für Leistungen wie Webhosting oder die Vermietung von Gameservern, doch diese wurden häufig nicht beglichen. Die daraus resultierenden Kosten für Inkassoverfahren, Mahn- und Vollstreckungsbescheide sowie gelegentliche Einsätze des freundlichen Gerichtsvollziehers waren enorm. Mein jetziges Geschäftsmodell ist jedoch grundlegend anders: Durch den Verkauf von lebenslangen Lizenzen als Einmalkauf entfällt das Problem der ausstehenden monatlichen Zahlungen.

Umfang von TAI

Thaliruth Arda Intelligence bietet bereits sehr viele Funktionen, was auch gewollt ist. Für meine persönlichen Anforderungen ist der Umfang ideal; was letztendlich du benötigst, liegt in deiner Entscheidung.

Abschlussbemerkung und Feedback für Sehende

Das Programm ist auch für sehende Menschen geeignet. Da ich selbst blind bin, liegt mein Fokus primär auf der Barrierefreiheit für blinde und sehbehinderte Nutzer. Daher kann ich die rein visuelle Gestaltung nicht abschließend beurteilen. Ich biete daher einer kleinen Gruppe sehender Tester eine kostenlose Demo-Lizenz an, um Feedback zur visuellen Qualität zu erhalten. So kann ich das Programm vor dem offiziellen Release noch einmal überarbeiten und visuelle Mängel beheben.

Diese Testphase ist natürlich mit Arbeitsaufwand verbunden. Wenn du mich dabei unterstützt, helfe ich dir später gerne bei einem Kauf mit einem Preisnachlass entgegen – eine Hand wäscht die andere.

Interessierte Sehende können mir gerne eine E-Mail an info@thaliruth.de schreiben.

Wenn du eine Demo-Lizenz erhältst, bitte ich dich um so detailliertes Feedback wie möglich. Ein Beispiel: Im Tab „Hörspiele“ ist die Schrift eher grau statt weiß und dadurch schwer lesbar. Für Screenreader-Nutzer ist dies kein Problem, für Sehende jedoch schon. Trotz präziser Anweisungen an die KI ist das Ergebnis nicht perfekt.

Melde dich gerne bei mir, denn Thaliruth Arda Intelligence wird definitiv mein Flaggschiff-Produkt werden.

1. Über dieses Handbuch

Dieses Handbuch beschreibt das Programm von innen: jeden Bereich, jede Schaltfläche, jede Einstellung. Es ist als Nachschlagewerk gedacht. Es liegt jeder Version in der neusten Fassung bei, aber auch als Dokumentation auf der Webseite <https://blindsoft.de> der.

Die beiden anderen Teile der Anleitung ergänzen dieses Handbuch: Der erste Teil beschreibt, was auf deinem Windows installiert sein muss, damit das Programm läuft. Der zweite Teil beschreibt, wie du die Zusatzfunktionen für Ton und Texterkennung einrichtest. Hier geht es allein um die Bedienung.

Ein Hinweis zur Schreibweise: Wenn im Handbuch von einem Bereich die Rede ist, ist damit einer der elf Hauptbereiche des Programmfensters gemeint, die du mit den Tastenkombinationen aus Kapitel 3 direkt anspringen kannst.

2. Der Aufbau des Programmfensters

Das Programm besteht aus einem einzigen Fenster, das beim Start immer maximiert erscheint. Darin gibt es elf Bereiche, zwischen denen du wechselst. Jeder Bereich hat eine eigene Tastenkombination, mit der du sofort dorthin springst, ohne dich durch eine Leiste tasten zu müssen.

Die elf Bereiche in ihrer Reihenfolge sind: Chat Übersicht, Kreatives, Audio Transkribierung, Text zu Sprache und Stimme klonen, Hörspiel und Podcast, Hörspiel-Werkstatt, Termine, Aufgaben, Ollama Modelle, KI-Verhalten und Einstellungen.

Wenn du in einen Bereich wechselst, sagt das Programm dessen Namen an und setzt den Fokus auf den Einleitungstext ganz oben. Von dort arbeitest du dich mit der TAB-Taste nach unten durch, oder du nutzt die Sprungtasten aus dem nächsten Kapitel, um schneller voranzukommen.

Ein Wort zu den Ansagen: Das Programm redet bewusst sparsam. Es meldet sich, wenn etwas fertig ist, wenn etwas schiefgeht oder wenn ein langer Vorgang noch läuft. Bei langen Vorgängen kommt in einem einstellbaren Takt eine Zwischenansage, damit du weißt, dass weitergearbeitet wird. Alles andere überlässt es deinem Screenreader.

3. Alle Tastenbefehle

Dieses Kapitel ist der wichtigste Teil des Handbuchs, denn über die Tastatur bist du im Programm am schnellsten. Die vollständige Liste bekommst du auch jederzeit im Programm selbst mit Strg plus Umschalt plus H angezeigt.

3.1. In einen Bereich wechseln

Strg + Umschalt + F1: Springt in die linke Seitenleiste auf die Überschrift Chats. Von dort geht es mit der TAB-Taste weiter zu den Projekten und zum Papierkorb.

Strg + Umschalt + F2: Zum Bereich Kreatives.

Strg + Umschalt + F3: Zum Bereich Audio Transkribierung.

Strg + Umschalt + F4: Zum Bereich Text zu Sprache und Stimme klonen.

Strg + Umschalt + F5: Zum Bereich Hörspiel und Podcast.

Strg + Umschalt + F6: Zum Bereich Hörspiel-Werkstatt.

Strg + Umschalt + F7: Zum Bereich Termine.

Strg + Umschalt + F8: Zum Bereich Aufgaben.

Strg + Umschalt + F10: Zum Bereich Ollama Modelle.

Strg + Umschalt + F11: Zum Bereich KI-Verhalten.

Strg + Umschalt + F12: Zum Bereich Einstellungen.

Die Taste F9 ist bewusst nicht belegt. Hier kommt vielleicht später eine weitere neue Funktion drauf.

3.2. Im Chat arbeiten

STRG + UMSCHALT + M: Springt direkt in das Eingabefeld. Wenn du dich bereits im Feld befindest, erscheint stattdessen der Hinweis, dass die Nachricht mit STRG + ENTER gesendet wird.

STRG + ENTER: Sendet die Nachricht ab. Die alleinige EINGABETASTE erzeugt im Eingabefeld lediglich eine neue Zeile. Ich habe mich bewusst für diese Steuerung entschieden, um zu verhindern, dass man beim schnellen Tippen oder beim Korrigieren von Fehlern versehentlich auf die ENTER-Taste kommt und so die Nachricht vorzeitig absendet. **Strg + Umschalt + N:** Legt einen neuen Chat an. Der Fokus springt danach ins Eingabefeld.

Strg + Umschalt + P: Legt ein neues Projekt an.

Strg + Umschalt + A: Springt auf die Modell-Auswahl.

Strg + Umschalt + U: Öffnet die Datei-Auswahl, um Dateien anzuhängen. Alternativ geht im Eingabefeld auch STRG + V, sofern eine Datei im Zwischenspeicher liegt.

Strg + Umschalt + Q: Stoppt eine laufende Antwort der KI. Solange etwas läuft, tut die Escape-Taste dasselbe.

Strg + Umschalt + V: Liest die letzte Antwort vor. Erneutes Drücken stoppt das Vorlesen.

Strg + Umschalt + K: Kopiert die letzte Antwort als reinen Text ohne Formatierungszeichen in die Zwischenablage.

Strg + Umschalt + X: Sagt die Zeichenanzahl der letzten Antwort an.

Strg + Umschalt + T: Sagt an, wie voll der Gesprächsspeicher des Modells ist, in Tokens und Prozent.

Strg + Alt + P: Springt in der Seitenleiste direkt auf den Papierkorb.

STRG + ALT + U: Erstellt für den aktuellen Chat ein Übergabeprotokoll inklusive Transkript. Dies ist besonders bei der Arbeit an Projekten von großem Vorteil, da eine KI dazu neigt, Inhalte zu vergessen oder zu

vermischen, sobald der Chatverlauf zu umfangreich wird. Ein „Token-Warner“ weist dich darauf hin, wenn dies der Fall ist – dazu später mehr. Mit diesem Übergabeprotokoll kannst du innerhalb eines Projekts einen neuen Chat starten und die KI direkt auf den aktuellen Stand bringen, damit du nahtlos dort weiterarbeiten kannst, wo du aufgehört hast.

3.3. Tastenbefehle, die überall gelten

Diese vier funktionieren auch dann, wenn das Programm gerade nicht im Vordergrund ist.

Strg + Umschalt + H: Zeigt die Übersicht aller Tastenbefehle.

Strg + Alt + T: Holt das Programm aus dem minimierten Zustand nach vorn und maximiert es.

Strg + Umschalt + Leertaste: Startet oder stoppt die Spracheingabe.

Strg + Umschalt + B: Macht ein Bildschirmfoto und sendet es an die KI.

3.4. Weitere Fenster öffnen

Strg + Umschalt + F: Öffnet das Fenster, mit dem du eine Rückmeldung an die Entwicklung senden kannst. Dazu später mehr.

Strg + Umschalt + I: Zeigt die Benachrichtigungen von Thaliruth Arda Intelligence an.

3.5. Sprungtasten innerhalb eines Bereichs

Weil die Schnellnavigation deines Screenreaders in dieser Art von Programm nicht greift, bringt das Programm eigene Sprungtasten mit. Sie beziehen sich immer auf den Bereich, in dem du dich gerade befindest. Mit zusätzlich gedrückter Umschalt-Taste geht es jeweils rückwärts.

Strg + Alt + H: Zur nächsten Überschrift.

Strg + Alt + B: Zur nächsten Schaltfläche.

Strg + Alt + F: Zum nächsten Eingabefeld.

Strg + Alt + L: Zum nächsten Link.

Dazu gibt es drei Übersichtslisten, aus denen du ein Ziel auswählst und direkt dorthin springst:

Strg + Alt + F5: Liste aller Eingabefelder des Bereichs.

Strg + Alt + F6: Liste aller Überschriften des Bereichs.

Strg + Alt + F7: Liste aller Links des Bereichs.

Im Chatverlauf gibt es außerdem Alt plus Pos1 und Alt plus Ende, um an den Anfang oder ans Ende des Gesprächs zu springen.

4. Der Bereich Chat Übersicht

Das ist dein Arbeitsplatz für das Gespräch mit der KI und zugleich der Ort, an dem du Chats und Projekte verwaltest. Du erreichst ihn mit Strg plus Umschalt plus F1.

4.1. Die linke Seitenleiste

Die TAB-Taste führt zuerst durch die gesamte Seitenleiste und danach erst in den Hauptbereich. Die Seitenleiste hat drei Abschnitte: Chats, Projekte und Papierkorb, jeweils mit einer eigenen Überschrift, die du direkt anspringen kannst.

Im Abschnitt Chats findest du zuerst ein Auswahlfeld mit dem Namen Filter Chats. Es hat vier Einträge: Aktive, Archiviert, Papierkorb und Alle. Der Filter bestimmt, welche Chats darunter erscheinen, und ändert zugleich den Text der Überschrift. Danach folgt die Schaltfläche Neuer Chat und darunter die Liste deiner Chats. Wenn du mit dem JAWS Fokus auf einem Chat stehst, der dir Namentlich angesagt wird, hast du folgende Funktionen:

Die Eingabetaste öffnet einen Chat, die Entf-Taste verschiebt ihn ohne Rückfrage in den Papierkorb, und die Kontextmenü-Taste öffnet weitere Befehle: Anpinnen, Umbenennen, Archivieren, In Projekt verschieben und Löschen. Ein angepinnter Chat wird mit dem Zusatz angepinnt angesagt.

Der Abschnitt Projekte ist ähnlich aufgebaut, mit einem eigenen Filter für aktive, archivierte oder alle Projekte und der Schaltfläche Neues Projekt. Jedes Projekt ist ein aufklappbarer Bereich: Mit der Eingabetaste oder der Leertaste klappt du es auf, und darunter erscheinen die Chats dieses Projekts. Im Kontextmenü eines Projekts kannst du Dateien anhängen,

die Anhänge verwalten, den Projekt-Kontext bearbeiten, einen neuen Chat im Projekt anlegen, das Projekt archivieren, umbenennen oder löschen.

Der Papierkorb ganz unten ist ebenfalls ein aufklappbarer Bereich, den du mit Strg plus Alt plus P direkt erreichst. Darin stehen gelöschte Chats und gelöschte Projekte getrennt unter eigenen Überschriften. Über das Kontextmenü stellst du einen Eintrag wieder her oder löschst ihn endgültig, wobei beim endgültigen Löschen eine Sicherheitsabfrage kommt. Die Schaltfläche Papierkorb leeren räumt alles auf einmal weg und nennt dir vorher die Anzahl der betroffenen Einträge.

4.2. Der Hauptbereich

Ganz oben steht die Modell-Auswahl, die du mit Strg plus Umschalt plus A direkt erreichst. Der erste Eintrag heißt Automatisch: Dann entscheidet das Programm anhand deiner Frage selbst, welches deiner Standardmodelle passt. Alle weiteren Einträge sind die bei dir installierten Modelle mit ihrem Originalnamen. Ist kein Modell verfügbar, steht an dieser Stelle ein Hinweis, dass du zuerst im Bereich Ollama Modelle eines installieren musst.

Darunter folgt der Chatverlauf. Jede Nachricht beginnt mit einer eigenen Überschrift. Bei deinen eigenen Nachrichten lautet sie Du, gefolgt von Wochentag, Datum und Uhrzeit. Bei Antworten der KI lautet sie Antwort von Thaliruth AI, gefolgt vom verwendeten Modell, Wochentag, Datum und Uhrzeit.

Zu jeder Antwort der KI gehören vier Schaltflächen: Vorlesen, Als Text kopieren, Mit Markdown kopieren und Zeichen zählen. Der Unterschied beim Kopieren ist wichtig: Als Text kopieren entfernt die Formatierungszeichen wie Rauten und Sternchen, behält aber Absätze und Listen, und ist damit für die meisten Zwecke die richtige Wahl. Mit Markdown kopieren übergibt den Rohtext, so wie das Modell ihn geschrieben hat. Code-Abschnitte in einer Antwort haben zusätzlich eine eigene Kopieren-Schaltfläche.

Unter dem Verlauf liegt das Eingabefeld. Die Eingabetaste erzeugt darin eine neue Zeile, gesendet wird mit Strg plus Eingabetaste. Wenn du im Eingabefeld Strg plus V drückst und in der Zwischenablage eine Datei

oder ein Bild liegt, wird es als Anhang übernommen statt eingefügt. Reiner Text wird ganz normal eingefügt.

Mit der Umschalt Taste + TAB, kannst du zu deinem Anhang springen und bei Bedarf löschen.

Ganz unten stehen die Schaltflächen Senden, Generierung stoppen und Dateien anhängen. Die Schaltfläche zum Stoppen ist nur bedienbar, solange wirklich eine Antwort erzeugt wird, sonst überspringst du sie beim Tabben.

4.3. Projekte: der eigene Arbeitsplatz für ein Vorhaben

Ein Projekt bündelt alles, was zu einem Vorhaben gehört: mehrere Chats, dauerhafte Hintergrund-Angaben, angehängte Dateien und, das ist der entscheidende Unterschied, einen echten Ordner auf deiner Festplatte. In einem Projekt-Chat darf die KI in diesem Ordner selbstständig Dateien anlegen, lesen, ändern und löschen. Außerhalb dieses Ordners kann sie das nicht, und zwar technisch nicht, nicht nur nach Absprache. Mehr dazu weiter unten.

Ein Projekt anlegen

Bevor du das erste Projekt anlegst, muss in den Einstellungen der Speicherpfad für Projekte gesetzt sein. Fehlt er, öffnet sich statt des Anlege-Fensters ein Hinweis, der dich zu dieser Einstellung führt. Alle Projektordner entstehen später unterhalb dieses Pfades.

Die Schaltfläche Neues Projekt in der Seitenleiste oder Strg plus Umschalt plus P öffnet das Anlege-Fenster mit vier Feldern:

Projektname: Der Name des Projekts. Er wird zugleich der Name des Ordners auf deiner Festplatte.

Projektart: Zwei Möglichkeiten: Coding und Programmierung oder Kreatives Projekt. Die Wahl entscheidet, welche Werkzeuge die KI in diesem Projekt bekommt, siehe unten.

Projekt-Kontext, optional: Ein mehrzeiliger Text mit dauerhaften Hintergrund-Angaben zu deinem Vorhaben. Er wird bei jeder Anfrage in diesem Projekt automatisch mitgeschickt, du musst ihn also nie wiederholen. Hier gehört hinein, was immer gelten soll, zum Beispiel

die verwendete Programmiersprache, der Erzählton einer Geschichte oder wiederkehrende Regeln.

Dateien anhängen, optional: Über die Schaltfläche Dateien auswählen hängst du gleich beim Anlegen Dateien an. In der Liste der ausgewählten Dateien bewegst du dich mit den Pfeiltasten, und die Entf-Taste nimmt eine Datei wieder heraus.

Mit der Schaltfläche Anlegen wird das Projekt erstellt und der Ordner auf der Festplatte angelegt.

Der Unterschied zwischen den beiden Projektarten

In beiden Projektarten darf die KI mit Dateien im Projektordner arbeiten. Der Unterschied liegt bei den Befehlen: Nur in einem Coding-Projekt bekommt die KI zusätzlich das Werkzeug, um Kommandozeilen-Befehle auszuführen, also zum Beispiel ein Programm zu bauen oder Tests laufen zu lassen. In einem kreativen Projekt gibt es dieses Werkzeug bewusst nicht.

Außerdem bekommt die KI in einem Coding-Projekt zusätzliche Verhaltensregeln mitgegeben, unter anderem zum Umgang mit Versionen und zum Löschen von Dateien.

Was die KI im Projektordner tun darf

Gehört ein Chat zu einem Projekt, stehen der KI sechs Datei-Werkzeuge zur Verfügung: eine Datei lesen, eine Datei schreiben, einen Ordner anlegen, eine Datei löschen, einen Ordner löschen und den Inhalt eines Ordners auflisten. Sie nutzt diese Werkzeuge selbstständig, wenn deine Aufgabe es erfordert. Bitte sie also ruhig direkt darum, ein Programm anzulegen oder eine Datei zu ändern; sie schreibt den Code dann nicht nur in den Chat, sondern legt ihn wirklich als Datei an. Bei jeder solchen Aktion hörst du eine kurze Ansage, zum Beispiel dass eine neue Datei angelegt wurde, samt Pfad.

Warum das sicher ist

Vor jedem Datei-Zugriff prüft das Programm den Zielpfad gegen den Projektordner. Diese Prüfung ist gründlich: Der Pfad wird zuerst bereinigt und vollständig aufgelöst, sodass Tricks mit zwei Punkten, die eine Ebene nach oben führen, ins Leere laufen. Danach wird geprüft, ob der aufgelöste Pfad wirklich innerhalb des Projektordners liegt. Ähnliche

Ordnernamen werden dabei sauber unterschieden, ein Projekt namens Mail bekommt also keinen Zugriff auf einen danebenliegenden Ordner namens Mail-Programm. Verweise auf andere Orte werden erkannt und abgelehnt, und die von Windows reservierten Gerätenamen sind als Dateinamen verboten.

Kurz gesagt: Die KI kann in deinem Projektordner alles tun und außerhalb davon nichts. Deine übrigen Dateien und dein System sind unberührt. In den Einstellungen kannst du zudem einen Backup-Pfad festlegen (Details dazu folgen weiter unten). Dies bewirkt, dass jedes Mal ein Backup der Projektdaten erstellt wird, wenn das entsprechende Modell eine Datei überarbeitet. So bleiben die Originaldatei sowie alle vorherigen Versionierungen stets erhalten. Es sollte jedoch angemerkt werden, dass das gesamte Projektsystem noch Fehler enthalten kann, da es bisher noch nicht umfassend getestet wurde. Ich werde jedoch mit kommenden Versionen entsprechende Verbesserungen nachreichen.

Befehle in Coding-Projekten

In einem Coding-Projekt darf die KI zusätzlich Kommandozeilen-Befehle ausführen, etwa um ein Programm zu bauen oder Tests laufen zu lassen. Dabei greifen zwei Sicherungen. Erstens gibt es in den Einstellungen die Liste der erlaubten Kommandozeilen-Befehle, ab Werk mit den gängigen Entwickler-Befehlen gefüllt. Zweitens wirst du vor jeder Ausführung gefragt und siehst den vollständigen Befehl; steht er nicht auf der Liste, bekommst du zusätzlich einen Warnhinweis. Ohne deine Bestätigung läuft nichts.

Projekt-Anhänge und Projekt-Kontext im Alltag

Beim ersten Prompt eines Chats bekommt die KI die Projekt-Anhänge vollständig mit. Bei den folgenden Nachrichten reicht ihr eine Liste der Dateinamen, und sie holt sich gezielt nach, was sie gerade braucht. Das spart Platz im Gespräch und Zeit. Über das Kontextmenü eines Projekts kannst du jederzeit Dateien anhängen, die Anhänge verwalten und den Projekt-Kontext bearbeiten.

Wenn eine Aufgabe zu groß wird

In den Einstellungen gibt es eine Obergrenze, wie viele Werkzeuge die KI in einer einzigen Antwort verwenden darf, ab Werk fünfundzwanzig. Legt

sie sehr viele Dateien an und erreicht diese Grenze, hängt das Programm an die Antwort einen Hinweis an, dass möglicherweise noch nicht alles fertig ist und du einfach mach weiter sagen kannst. Die KI sieht dann nach, was bereits angelegt wurde, und arbeitet weiter. Bei großen Vorhaben lohnt es sich, diese Grenze in den Einstellungen zu erhöhen.

Sicherungskopien

Schaltest du in den Einstellungen die automatischen Backups ein, legt das Programm eine Sicherungskopie an, bevor die KI eine vorhandene Datei ändert oder löscht. Bei Projekten, an denen dir viel liegt, ist das eine sinnvolle Vorsichtsmaßnahme.

4.4. Wenn aus einer Antwort ein fertiges Dokument wird

Eine der praktischsten Eigenschaften des Programms: Die KI antwortet nicht nur im Chat, sie kann das Ergebnis auch direkt als fertiges Dokument anlegen und öffnen. Das funktioniert in jedem Chat, auch ohne Projekt.

Ein Beispiel aus dem Alltag

Du bittest die KI um ein Rezept für einen Apfelkuchen. Sie antwortet im Chat und legt zugleich ein sauber gegliedertes Word-Dokument mit dem vollständigen Rezept an, das sich anschließend von selbst öffnet. Du hörst dabei eine Ansage mit dem Dateinamen. Damit kannst du das Rezept sofort weiterbearbeiten, ausdrucken oder dauerhaft ablegen, ohne etwas aus dem Chat herauskopieren zu müssen.

Die drei Dokumentarten

Die KI entscheidet selbst, welche Art zur Aufgabe passt, und du kannst sie auch ausdrücklich darum bitten.

Word-Dokument: Ein echtes, barrierefrei gegliedertes Dokument mit Überschriften, Absätzen und Listen. Der Titel wird zur Überschrift der obersten Ebene und zugleich zum Dateinamen. Das ist die richtige Wahl für alles, was du weiterbearbeiten möchtest.

PDF-Dokument: Ebenfalls gegliedert, mit echtem, auswählbarem Text. Die richtige Wahl zum Weitergeben und Archivieren.

Textdatei: Reiner Text ohne Formatierung. Für Notizen und für alles, was schlicht bleiben soll.

Eingescannte PDF-Dateien lesbar machen

Hängst du ein eingescanntes PDF an, also eines, das nur aus Bildern besteht und deshalb für deinen Screenreader stumm ist, bekommt die KI ein zusätzliches Werkzeug angeboten. Damit erkennt das Programm den Text und baut daraus ein sauber gegliedertes PDF mit echtem Text. Wichtig dabei: Bei diesem Weg formuliert die KI nichts um. Sie liefert keinen Inhalt, sondern der erkannte Text wird wortgetreu übernommen. Dieses Werkzeug erscheint nur, wenn wirklich ein eingescanntes PDF angehängt ist. Voraussetzung ist die eingerichtete Texterkennung, siehe Teil 2 der Anleitung.

Wo die Dateien landen

Gehört der Chat zu einem Projekt, landet das Dokument im Projektordner. In einem Chat ohne Projekt landet es in dem Ordner, den du in den Einstellungen unter Speicherpfad für Dokumente eingetragen hast; ist das Feld leer, ist es dein Downloads-Ordner. Gibt es dort schon eine Datei desselben Namens, wird durchnummeriert, es geht also nichts verloren.

Ausdrücklich ein Word-Dokument verlangen

Bei sehr langen Inhalten, etwa der Übersetzung einer ganzen Datei, gibt es einen zweiten, besseren Weg. Formulierst du in deiner Nachricht ausdrücklich den Wunsch nach einem Word-Dokument, erkennt das Programm das und geht anders vor: Es lässt die KI ganz normal antworten und baut das Dokument anschließend selbst aus der fertigen Antwort. Der Vorteil ist, dass so beliebig lange Inhalte funktionieren, während der übliche Weg bei sehr großen Mengen an eine Zeitgrenze stoßen kann.

Erkannt wird dieser Wunsch an zwei Dingen zusammen: Deine Nachricht enthält ein Wort wie Word oder docx, und sie enthält ein Tätigkeitswort wie speichern, erstellen, erzeugen, schreiben, anlegen, exportieren oder umwandeln. Ein Beispiel: Übersetze diese Datei und speichere sie als Word-Dokument. Ausgenommen sind Fragen: Endet deine Nachricht mit einem Fragezeichen, löst sie kein Dokument aus. Die Frage, ob das Programm Word-Dokumente erstellen kann, wird also beantwortet und nicht ausgeführt.

4.5. Was beim Senden im Hintergrund passiert

Anhänge werden je nach Art unterschiedlich behandelt. Bilder gehen an ein bildfähiges Modell. Bei Textdateien wird der Inhalt ausgelesen und der Frage beigelegt. Bei Dateien, die das Programm noch nicht auswerten kann, etwa ZIP, wird die Datei zwar angehängt und gespeichert, das Modell bekommt aber den ehrlichen Hinweis, dass der Inhalt noch nicht lesbar ist.

Enthält deine Nachricht eines der Auslöser-Wörter wie suche, recherchiere, nachschlagen oder im Netz, sucht das Programm im Internet. Unabhängig davon kann das Modell auch selbst entscheiden, dass eine Suche nötig ist. Ist in den Einstellungen die Nachfrage eingeschaltet, wirst du vorher gefragt und siehst die Suchbegriffe.

Steht in deiner Nachricht eine Internetadresse und ist der Webseiten-Abruf eingeschaltet, holt das Programm die erste Adresse selbst ab und legt den Seitentext der Frage bei. Du hörst dabei die Ansage, dass die Seite abgerufen wird.

Beginnt deine Nachricht mit merke dir, speichert das Programm den Rest als Erinnerung. Zu jeder Frage sucht es außerdem selbst die passenden gespeicherten Erinnerungen heraus und gibt sie dem Modell mit.

Die KI kann darüber hinaus Werkzeuge benutzen: im Web suchen, Aufgaben und Termine anlegen, Erinnerungen speichern und lesen sowie Dokumente erzeugen. Wie aus einer Antwort ein fertiges Word-Dokument, ein PDF oder eine Textdatei wird, steht ausführlich im Abschnitt 4.4. Die Werkzeuge zum Anlegen und Ändern von Dateien und zum Ausführen von Befehlen stehen nur in Projekt-Chats zur Verfügung; alles dazu steht im Abschnitt 4.3.

5. Der Bereich Kreatives

Hier schreibst du längere erzählende Texte mit Hilfe eines Sprachmodells. Du erreichst den Bereich mit Strg plus Umschalt plus F2, Er gehört zum Premium-Umfang, alternativ ist diese Funktion auch einzeln zubuchbar. und ist nur bedienbar, wenn deine Lizenz ihn enthält.

5.1. Wofür der Bereich gedacht ist

Du kannst zwei Dinge tun: eine neue Geschichte aus einer Beschreibung anlegen oder eine vorhandene Geschichte um einen weiteren Teil fortsetzen. Jeder fertige Teil wird als barrierefreies Word-Dokument gespeichert und danach automatisch geöffnet. Eine Geschichte kann eine fortlaufende Reihe aus mehreren Teilen sein oder eine einzelne, abgeschlossene Erzählung.

5.2. Die Bedienelemente von oben nach unten

Ganz oben wählst du das Modell, mit dem geschrieben wird. Anders als im Chat gibt es hier bewusst keinen Eintrag Automatisch. Vorbelegt ist das zuletzt gewählte Modell, sonst dein Standardmodell für kreative Texte.

Darunter steht das Feld Ziel-Zeichenzahl pro Teil. Der Vorgabewert ist 5000 Zeichen, erlaubt sind 500 bis 200000. Das ist ein Ziel, keine harte Grenze.

Im Abschnitt Neue Geschichte anlegen beschreibst du im großen Feld, worum es gehen soll und womit die Geschichte beginnt. Das Kontrollfeld Fortlaufende Reihe ist ab Werk angehakt und bedeutet eine Reihe aus mehreren Teilen; nimmst du das Häkchen weg, entsteht eine einzelne abgeschlossene Geschichte. Das Feld Titel darfst du leer lassen, dann schlägt das Modell selbst einen vor. Die Schaltfläche Neue Geschichte jetzt erzeugen startet die Arbeit.

Im Abschnitt Geschichte fortsetzen wählst du aus der Liste eine vorhandene Geschichte. Jeder Eintrag wird mit dem Haupttitel und dem Titel des letzten Teils angesagt. Im Feld Handlungsvorgabe kannst du beschreiben, wie der nächste Teil beginnen soll; lässt du es leer, entscheidet das Modell selbst, wie es weitergeht. Auch hier ist der Titel des neuen Teils freiwillig.

Ganz unten steht der Statusbereich, der dir jederzeit sagt, was gerade passiert.

5.3. Was du während der Arbeit hörst

Beide Schaltflächen bleiben während des Schreibens bedienbar, damit dein Fokus nicht verloren geht. Drückst du versehentlich ein zweites Mal, hörst du nur den Hinweis, dass bereits eine Erzeugung läuft. In deinem eingestellten Takt sagt das Programm den Stand an, zum Beispiel dass die Geschichte vorbereitet wird, wie viele Zeichen geschrieben sind, dass

der Abschnitt geprüft wird oder dass die Geschichte fertig ist und gespeichert wird.

Am Ende hörst du eine ausführliche Erfolgsansage mit dem Titel, dem Dateinamen und dem Ordner. Wurde der Titel vom Modell vorgeschlagen, wird das ausdrücklich dazugesagt. Danach öffnet sich das Dokument. Abbrechen kannst du jederzeit mit Escape oder Strg plus Umschalt plus Q.

6. Der Bereich Audio Transkribierung

Hier wandelst du eine Tonaufnahme in Text um. Du erreichst den Bereich mit Strg plus Umschalt plus F3, auch er ist Teil des Premium Umfangs, und kann auch einzeln dazu gebucht werden.

Wichtig vorab: Diese Funktion braucht derzeit zwingend eine NVIDIA-Grafikkarte und ein von Hand installiertes WhisperX. Näheres steht im zweiten Teil der Anleitung.

6.1. Die Aufnahme laden

Es gibt drei Wege zur Aufnahme. Erstens die Schaltfläche Audio-Datei durchsuchen, die den Windows-Dateidialog öffnet. Zweitens Strg plus V, wenn du die Datei zuvor im Explorer kopiert hast. Drittens der Weg über YouTube. Unterstützt werden die Formate mp3, m4a, wav, ogg, flac, aac und wma. Unter der Schaltfläche steht immer, welche Datei gerade geladen ist.

Für den YouTube-Weg fügst du die Adresse in das Feld YouTube-Adresse ein und drückst die Schaltfläche Video-Ton laden und transkribieren. Das Programm lädt dann die Tonspur als MP3, macht sie zur geladenen Datei und startet unmittelbar danach die Transkription mit deinen Einstellungen. Die MP3 bleibt im Ordner für Übergaben und Transkripte liegen.

6.2. Die Optionen

Automatisch ins Deutsche übersetzen ist ab Werk angehakt und wirkt nur, wenn die Aufnahme als englisch erkannt wird.

Sprechertrennung ist ebenfalls ab Werk angehakt. Wenn du sie einschaltest, erscheinen sofort ein Feld für die Anzahl der Sprecher und acht Felder für die Namen. Die Anzahl darfst du leer lassen, dann entscheidet die Trennung selbst; gibst du eine Zahl von eins bis acht an, ist das die stärkste Hilfe für ein sauberes Ergebnis. Die Namensfelder sind freiwillig; bleibt eines leer, heißt die Person im Text einfach Sprecher mit der jeweiligen Nummer.

Zeitstempel ist ab Werk nicht angehakt. Schaltest du es ein, stehen im fertigen Text Zeitmarken.

Alle Optionen und die Sprecheramen werden sofort gespeichert und stehen beim nächsten Mal wieder so da.

6.3. Ein wichtiger Hinweis zur Sprechertrennung

Wenn die Sprechertrennung angehakt ist, in den Einstellungen aber kein Hugging-Face-Schlüssel hinterlegt wurde, überspringt das Programm die Trennung stillschweigend. Es gibt dazu keine Ansage, sondern nur einen Eintrag im Protokoll. Wenn du dich also wunderst, warum im fertigen Text keine Sprecher unterschieden werden, prüfe zuerst den Schlüssel in den Einstellungen.

6.4. Das Ergebnis

Nach dem Start hörst du in deinem eingestellten Takt Zwischenansagen. Am Ende meldet das Programm, unter welchem Namen gespeichert wurde, ob übersetzt wurde und wo das englische Original liegt, und öffnet das Word-Dokument. Alles landet im Ordner für Übergaben und Transkripte. Abbrechen geht jederzeit mit Escape oder Strg plus Umschalt plus Q. Als Notbremse beendet das Programm einen Lauf, der länger als sechs Stunden dauert.

7. Der Bereich Text zu Sprache und Stimme klonen

Hier lässt du geschriebenen Text in einer natürlichen Stimme sprechen. Du erreichst den Bereich mit Strg plus Umschalt plus F4, er ist auch Teil des Premium Umfangs, kann aber auch einzeln dazu gebucht werden, und setzt die eingerichtete Chatterbox-Umgebung voraus.

7.1. Der Ablauf

Du gibst deinen Text in das große Feld ganz oben ein. Darunter wählst du eine Stimme, stellst die drei Regler ein und drückst die Schaltfläche KI-Stimme erzeugen und abspielen. Die Aufnahme wird sofort abgespielt. Gefällt sie dir, trägst du unten einen Dateinamen ein und speicherst sie.

7.2. Die Stimmen

Die Stimmen kommen aus dem Unterordner Referenzstimmen deines Stimmen-Speicherpfads. Jede WAV-Datei dort wird zu einem Eintrag in der Auswahlliste, benannt nach dem Dateinamen ohne Endung. Eine Referenz-Aufnahme von etwa zehn bis zwanzig Sekunden genügt. Genau das ist das Stimmen-Klonen: Die erzeugte Sprache übernimmt Klang und Sprechweise deiner Vorlage. Deshalb gibt es keinen getrennten Klon-Bereich. Der Ordner Referenz-Stimmen muss manuell angelegt werden! Du musst also zunächst selbst eine Stimme aufzeichnen, um Texte später in diese umwandeln zu können. Nutze dazu beispielsweise den Windows Audio Recorder, nimm einige Sekunden auf und speichere die fertige Datei im Ordner Referenz-Stimmen.

(Wichtiger Hinweis für Screenreader-Nutzer: Der Ordner muss exakt als Referenz-Stimmen geschrieben werden – es muss ein Bindestrich zwischen „Referenz“ und „Stimmen“ stehen und es dürfen keine Leerzeichen verwendet werden.)

Zuvor musst du unter den Einstellungen bei den Speicherpfaden einen Ort für die geklonten Stimmen festlegen, zum Beispiel: D:\Thaliruth-AI\Voices. In diesem Ordner Voices werden alle generierten Stimmen gespeichert. Innerhalb dieses Verzeichnisses muss sich der Ordner Referenz-Stimmen befinden, in den du deine Aufnahme legst, damit das System ordnungsgemäß funktioniert.

Solange dort noch keine Datei liegt, siehst du einen Hinweis und die Auswahlliste bleibt leer. Nachdem du eine Aufnahme abgelegt hast, drückst du die Schaltfläche Stimmen aktualisieren; das Programm liest den Ordner dann neu ein und sagt an, wie viele Stimmen es gefunden hat. Mit der Schaltfläche Ausgewählte Stimme als Standard festlegen merkt sich das Programm deine bevorzugte Stimme dauerhaft. Die

Schaltfläche Probe anhören spricht einen kurzen festen Satz, damit du eine Stimme schnell beurteilen kannst, ohne eigenen Text einzugeben.

7.3. Die drei Regler

Alle drei sind bewusst normale Eingabefelder für Zahlen und keine Schieberegler. Der Grund: Bei einem Schieberegler liest der Screenreader die Stellung als Prozentwert vor, in einem Eingabefeld hörst du dagegen den tatsächlichen Wert. Die Felder nehmen nur Ziffern an, und beim Verlassen wird ein zu großer oder zu kleiner Wert zurechtgerückt und die Korrektur angesagt.

Emotionalität: Stufe 0 bis 10, Vorgabe 5. Null klingt neutral und sachlich, zehn stark dramatisch.

Sprechtempo: Stufe 0 bis 10, Vorgabe 5. Niedrige Werte sprechen langsam und bewusst, hohe Werte zügig.

Wiederholungsbremse: Stufe 0 bis 10, Vorgabe 10. Sie verhindert das kurze, abgehackte Wiederholen am Satzende, das vor allem bei sehr kurzen Texten auftreten kann. Der hohe Vorgabewert ist Absicht.

7.4. Mitten im Text die Stimme wechseln

Du kannst innerhalb eines Textes zwischen mehreren Stimmen wechseln. Dazu setzt du an der Stelle, an der die neue Stimme beginnen soll, eine Markierung: eine eckige Klammer auf, dann der Name der Referenzstimme ohne die Endung, dann eine eckige Klammer zu. Der Text vor der ersten Markierung wird mit der oben ausgewählten Stimme gesprochen. Ab einer Markierung gilt die dort genannte Stimme, bis die nächste folgt. Die Markierungen selbst werden nicht mitgesprochen.

Verweist eine Markierung auf eine Stimme, die es nicht gibt, sagt das Programm dir das mit Namen, damit du entweder die Aufnahme nachlegen oder die Markierung berichtigen kannst.

7.5. Speichern

Im Feld Dateiname zum Speichern trägst du einen Namen ohne Endung ein; unzulässige Zeichen ersetzt das Programm selbst. Lässt du das Feld leer, entsteht ein Name aus dem Wort KI-Stimme mit Datum und Uhrzeit. Vorhandene Dateien werden nie überschrieben, stattdessen wird

durchnummeriert. Die Schaltfläche Speicherort öffnen zeigt dir den Ordner im Explorer.

8. Der Bereich Hörspiel und Podcast

Hier entsteht aus einer fertig gesprochenen Aufnahme ein vollständiges Hörspiel mit Geräuschen und Musik. Du erreichst den Bereich mit Strg plus Umschalt plus F5. Er ist auch Teil des Premium Umfangs, kann aber auch einzeln dazu gebucht werden, und braucht die eingerichtete Klang-Umgebung sowie eine NVIDIA-Grafikkarte.

8.1. Der Ablauf in Stufen

Dieser Bereich arbeitet in Stufen, die aufeinander aufbauen. Jede Stufe schreibt ihr Ergebnis als Datei in einen Arbeitsordner, den die nächste Stufe liest. Die Reihenfolge ist verbindlich: Überspringst du eine Stufe, sagt dir das Programm freundlich, was zuerst fehlt.

1. Aufnahme laden. Über die Schaltfläche Audio-Datei durchsuchen oder mit Strg plus V.
2. Originaltext einfügen. Entweder mit Text aus Datei laden aus einer Text-, Word- oder PDF-Datei, oder direkt in das große Feld tippen. Hier stehen auch deine eigenen Markierungen. Mit „Markierungen“ ist gemeint, dass du hier eine Stimme oder auch Soundeffekte auswählen kannst. Ein kleines Beispiel: Eckige Klammer auf, [Thaliruth] Eckige Klammer zu, Hallo, ich bin Thaliruth, und Eckige Klammer auf, [furz] Eckige Klammer zu, furze gern. Die erste Markierung wählt die Stimme aus, mit der dieses Hörspiel gesprochen wird. Die zweite Markierung fügt einen kleinen Soundeffekt hinzu. Dieser wird dann genau bei dem Wort Furz in diesem Beispiel eingebaut. Dieser muss natürlich zuvor in der Werkstatt erzeugt und entsprechend auf deinem System im Ordner „Klänge“ abgespeichert worden sein, damit er am Ende korrekt in das Hörspiel eingebaut werden kann.
3. Zeitbasis erstellen. Das Programm hört die Aufnahme ab und hält fest, welches Wort zu welcher Sekunde fällt. Erst dieser Schritt legt den Projektordner an. Danach hörst du die Wortanzahl und Gesamtdauer.

4. Klangliste erstellen. Ein Sprachmodell liest deinen Text und sucht Geräusche und die passende Musikstimmung heraus. Danach hörst du, wie viele Klangereignisse gefunden wurden.
5. Geräusche erzeugen. Aus der Klangliste entstehen die einzelnen Tondateien.
6. Musik einstellen und erzeugen. Die Musik entsteht mit Absicht als eigene Datei und wird noch nicht abgespielt, damit du sie zuerst pur anhören kannst.
7. Hörspiel mischen. Der letzte Schritt legt Geräusche und Musik unter die Erzählerstimme. Hier rechnet keine KI, das dauert nur Sekunden. Danach wird das Ergebnis abgespielt (Platzhalter: hier wird später ein Link zum Podcast Eingebunden, wo du dir dann ein Hörspieldemo mit Sprecherstimme und Soundeffekten anhören kannst.).
8. Ein Projekt später weiterbearbeiten

Ganz oben im Bereich steht das Auswahlfeld Projekt zum Überarbeiten. Damit lädst du ein früheres Hörspiel zurück, samt Text, Intro- und Outro-Texten und der Aufnahme. Das Programm sagt dir dabei, ob schon eine Zeitbasis vorhanden ist und du gleich bei der Klangliste weitermachen kannst. Danach lässt du gezielt einzelne Stufen neu laufen. Die Schaltfläche Projektliste neu laden liest die Projektordner erneut ein.

8.3. Die Musik einstellen

Es gibt drei Arten von Musik, die du einzeln oder gemeinsam anhängen kannst.

Intro und Outro: Ein Vorspann vor der ersten Zeile und ein Nachspann nach der letzten. Bei langen Geschichten ist das oft die bessere Wahl als durchgehende Musik.

Durchgehende Musik: Ein Musikbett unter der ganzen Erzählung. Das kostet viel Rechenzeit, denn das Bett entsteht aus Kacheln von je drei Minuten, die unhörbar ineinander übergehen. Für eine Stunde Erzählung sind das einundzwanzig Stücke.

Kurze Spannungsmusik: Kurze Stücke an einzelnen Stellen, dort wo die Spannung steigt. So arbeiten echte Hörspiele. Dafür liest das Sprachmodell deinen Text ein zweites Mal, diesmal auf der Suche nach

Dramaturgie. Zwischen zwei Stellen liegen mindestens neunzig Sekunden.

Die Musikstimmung trägst du in englischen Stilbegriffen ein, zum Beispiel dark medieval, celtic folk oder epic orchestral. Das Modell versteht Englisch am besten. Hat die Klangliste bereits einen Vorschlag geliefert und ist das Feld noch leer, wird der Vorschlag automatisch eingetragen; deine eigene Eingabe hat aber immer Vorrang.

Statt Musik erzeugen zu lassen, kannst du für Intro und Outro auch je eine eigene Datei hinterlegen. Dafür gibt es je eine Schaltfläche zum Wählen und eine zum Entfernen. Zusätzlich kannst du einen Intro-Text und einen Outro-Text eingeben, die dann mit deiner Stimme gesprochen werden.

Das Auswahlfeld Ausprägung des Musikmodells hat drei Einträge. Beste Qualität rechnet fünfzig Schritte und klingt hörbar besser, dauert aber etwa sechsmal so lange. Schnell rechnet acht Schritte und eignet sich, um in wenigen Minuten zu prüfen, ob überhaupt alles läuft. Der dritte Eintrag ist das Grundmodell.

8.4. Die Mischung

Vor dem Mischen stellst du zwei Lautstärken ein, jeweils als Stufe von 0 bis 30. Stufe 30 bedeutet so laut wie die Erzählerstimme. Für die Geräusche ist 30 voreingestellt, für die Musik 12, was deutlich leiser ist. Das Kontrollfeld Musik mit einmischen ist voreingestellt. Die Musik wird während der Mischung automatisch geregelt: Unter der Sprache tritt sie zurück, in den Sprechpausen kommt sie von selbst hoch.

8.5. Der Podcast-Weg

Ganz unten im Bereich gibt es einen eigenen, viel kürzeren Weg für Podcast-Folgen. Deine Aufnahme bleibt dabei völlig unverändert, es kommen nur Intro und Outro dazu. Zeitbasis, Klangliste und Geräusche brauchst du dafür nicht. Du klickst oben im Abschnitt Musik das Kontrollfeld Intro und Outro an, stellst die Stimmung ein oder hinterlegst eigene Dateien, wählst mit Podcast-Aufnahme wählen deine Datei und drückst Podcast fertigstellen.

8.6. Wo die Dateien landen

Für jedes Hörspiel entsteht ein eigener Projektordner, benannt nach deiner Aufnahme. Darin liegen unter anderem der Originaltext, die Zeitbasis, die Klangliste, ein Unterordner mit den erzeugten Geräuschen, ein Unterordner mit der Musik und am Ende die fertige Datei. Das fertige Hörspiel heißt fertiges-hörspiel.mp3, das Ergebnis des Podcast-Wegs heißt fertiger-podcast.mp3. Eigene Klänge und eigene Musik liegen nicht im Projektordner, sondern in den beiden Ordnern, die du in den Einstellungen festgelegt hast.

Alte Hörspiele räumt das Programm beim Start selbst auf. Aufräumen heißt dabei ausdrücklich nicht löschen: Nur die großen Zwischendateien verschwinden, das fertige Hörspiel und dein Originaltext bleiben erhalten. Die fünf neuesten Ordner werden nie angetastet.

8.7. Eigene Klänge und Musik einbauen

Du musst nicht alles erzeugen lassen. Mit einer Markierung im Text baust du eigene Dateien ein: eckige Klammer auf, das Wort klang, ein Doppelpunkt, der Dateiname, eckige Klammer zu. Für Musik lautet das Wort musik statt klang. Das Programm sucht die Datei dann in deinem Ordner für eigene Klänge beziehungsweise eigene Musik. Stehen in der Klangliste ausschließlich eigene Klänge, wird gar kein Modell geladen und der Schritt geht sehr schnell.

9. Der Bereich Hörspiel-Werkstatt

Die Werkstatt erzeugt einzelne Geräusche und einzelne Musikstücke, ohne dass du ein Hörspiel anlegen musst. Du erreichst sie mit Strg plus Umschalt plus F6, diese Werkstatt ist auch Teil des Premium Umfangs, kann aber auch einzeln dazu gebucht werden.

9.1. Der Gedanke dahinter

Nicht jeder Versuch des Modells sitzt. Deshalb erzeugt ein Lauf gleich mehrere Varianten zum selben Begriff. Du hörst sie nacheinander an, speicherst die guten und verwirfst den Rest. Das Modell wird dabei nur einmal geladen, was den Löwenanteil der Zeit ausmacht; mehrere Einzelaufrufe wären deutlich langsamer.

Was du speicherst, landet genau in den Ordnern, aus denen die Markierungen im Hörspiel ihre Dateien holen. Damit schließt sich der

Kreis: In der Werkstatt baust du dir deinen eigenen Vorrat auf und rufst ihn später im Text mit einer Markierung ab.

Wichtig: Ein neuer Lauf verwirft die bisherigen Varianten. Speichere also vorher, was du behalten möchtest.

9.2. Der Aufbau

Die Werkstatt hat zwei getrennte Abschnitte, Geräusch erzeugen und Musik erzeugen. Beide Abschnitte findest du schnell mit Strg plus Alt plus H.

Beide sind gleich aufgebaut: ein Feld für den Begriff auf Englisch, Felder für Länge und Anzahl der Varianten, eine Schaltfläche zum Erzeugen, ein Feld für den Namen zum Speichern und darunter die Liste der Varianten. Zu jeder Variante gehören drei Schaltflächen: Anhören, Speichern und Verwerfen. Sie werden eindeutig angesagt, zum Beispiel Variante eins anhören. Am Ende jedes Abschnitts steht eine Schaltfläche, mit der du alle Varianten auf einmal verwirfst.

9.3. Geräusche

Die Länge reicht von 3 bis 20 Sekunden, voreingestellt sind 6, Die Anzahl der Varianten reicht von 1 bis 8, voreingestellt sind vier. Ein ehrlicher Hinweis aus dem Programm: Alltägliche Klangquellen kann das Modell gut, also Wind, Schritte oder ein Knurren. Fantasie-Klänge kann es nicht. Ein Vergleich aus dem Alltag bringt darum mehr als eine Beschreibung von Magie.

9.4. Musik

Hier gibt es zusätzlich das Auswahlfeld Stimmung mit sechzehn vorbereiteten Stilen, von Gasthaus warm und gemütlich über Düsterwald unheimlich bis Elbengesang ätherisch. Diese Voreinstellungen sind für meine Verwendungen eingebaut worden. Ich kann diese Liste später aber gegen eure Wünsche anpassen und meine Stile entfernen. Die Auswahl füllt das Begriff-Feld mit einem passenden englischen Stil, den du danach frei ändern kannst. Der oberste Eintrag Eigener Stil ändert nichts.

Das Feld Gesang, Silben oder Text ist freiwillig. Lässt du es leer, entsteht Musik ohne Gesang. Für wortloses Trällern eignen sich Silben, für einen Titelsong ganze Zeilen. Die Länge reicht von 10 bis 240 Sekunden,

voreingestellt sind 30. Die Musik läuft in der Werkstatt immer in der besten Qualität, hier gibt es keine Modellwahl.

9.5. Speichern

Trage den Namen am besten ein, bevor du speicherst. Ist der Name schon vergeben, zählt das Programm hoch und nennt dir in der Ansage den tatsächlich verwendeten Namen, gleich zusammen mit der Markierung, mit der du die Datei später im Text aufrufst. Eine vorhandene Datei wird niemals überschrieben.

10. Der Bereich Termine

Hier verwaltest du Termine mit Datum und Uhrzeit. Du erreichst den Bereich mit Strg plus Umschalt plus F7, er ist auch Teil des Premium Umfang, kann aber auch einzeln dazu gebucht werden.

10.1. Termine anlegen und bearbeiten

Die Schaltfläche Neuen Termin anlegen öffnet ein Fenster mit fünf Feldern: Bezeichnung, Datum, Uhrzeit, Wiederholung und eine freiwillige Notiz. Datum und Uhrzeit sind bewusst einfache Textfelder statt eines Kalender-Steuerelements. Das Datum gibst du als Tag Punkt Monat Punkt Jahr ein, die Uhrzeit als Stunde Doppelpunkt Minute. Bei einem neuen Termin sind das heutige Datum und acht Uhr vorbelegt, und der Fokus steht im Feld Bezeichnung.

Die Wiederholung hat fünf Möglichkeiten: einmalig, täglich, wöchentlich, monatlich und jährlich. Bestätigt wird mit der Eingabetaste, abgebrochen mit Escape. Ist eine Eingabe nicht lesbar, bleibt das Fenster offen und der Fokus springt auf das betroffene Feld.

10.2. Die Terminliste

Die Termine stehen nach Fälligkeit sortiert, der nächste oben. Jede Zeile nennt Bezeichnung, Datum und Uhrzeit, bei wiederkehrenden Terminen die Wiederholungsart, bei vorhandener Notiz den Zusatz mit Notiz und bei abgelaufener Zeit den Zusatz überfällig.

Auf einer Zeile öffnet die Eingabetaste die Bearbeitung, die Entf-Taste löscht den Termin sofort ohne Rückfrage, und die Kontextmenü-Taste bietet Bearbeiten, Notiz anzeigen, Als erledigt markieren und Löschen an.

Notizen erscheinen in einem eigenen kleinen Fenster, damit dein Screenreader sie zuverlässig vorliest.

Beim Erledigen gibt es einen wichtigen Unterschied: Ein einmaliger Termin verschwindet, ein wiederkehrender rückt um genau ein Intervall weiter, und das Programm sagt dir den neuen Termin gleich an.

10.3. Die Termin-Erinnerung

Solange das Programm läuft, prüft es alle dreißig Sekunden, ob ein Termin fällig ist. Ist es soweit, spielt es einen Ton ab und öffnet ein Erinnerungsfenster. War das Programmfenster minimiert, holt es sich vorher selbst in den Vordergrund, denn sonst könnte dein Screenreader das Fenster nicht erfassen.

Im Erinnerungsfenster steht je fälligem Termin ein Kontrollfeld, und der Fokus liegt gleich auf dem ersten davon. Du kreuzt an, was du erledigt hast. Darunter gibst du an, in wie vielen Minuten an den Rest erneut erinnert werden soll, erlaubt sind 1 bis 120, voreingestellt sind 5. Die Schaltfläche Bestätigen markiert die angekreuzten Termine als erledigt und verschiebt den Rest. Später erinnern verschiebt alles. Das Schließen über das Fensterkreuz zählt wie Später erinnern.

Eine Besonderheit: Termine, die schon vor dem Programmstart überfällig waren, lösen keine Erinnerung aus. Sie erscheinen stattdessen im Assistenz-Fenster beim Start, siehe Kapitel 15.

10.4. Zusammenspiel mit der KI

Die KI kann auf deine Bitte hin Termine anlegen, wenn in den Einstellungen der Schalter KI darf Aufgaben und Termine anlegen eingeschaltet ist. Relative Angaben wie morgen rechnet sie selbst um. Gibt es zur selben Uhrzeit schon einen Termin, fragt dich das Programm in einem eigenen Fenster, ob du trotzdem eintragen möchtest; der Fokus steht dabei auf Nicht eintragen. Lehnst du ab, bekommt auch die KI die klare Rückmeldung, damit sie nicht behauptet, sie hätte eingetragen.

Termine abfragen kann die KI bewusst nicht. Dafür genügt das Assistenz-Fenster beim Programmstart.

11. Der Bereich Aufgaben

Hier führst du eine dauerhafte Aufgabenliste ohne Zeitbezug. Du erreichst den Bereich mit Strg plus Umschalt plus F8, er ist auch Teil des Premium Umfangs, kann aber auch einzeln dazu gebucht werden.

11.1. Aufgaben anlegen

Ganz oben steht ein Eingabefeld für den Aufgabentext. Die Eingabetaste in diesem Feld legt die Aufgabe direkt an, du musst also nicht erst zur Schaltfläche Aufgabe hinzufügen tabben. Danach ist das Feld wieder leer und bereit für die nächste Aufgabe.

11.2. Die Aufgabenliste

Aufgaben haben genau zwei Ebenen: Hauptaufgaben und Unteraufgaben. Die Unteraufgaben stehen direkt unter ihrer Hauptaufgabe. Jede Zeile nennt bei einer Unteraufgabe zuerst das Wort Unteraufgabe, dann den Titel und den Zustand erledigt oder nicht erledigt. Bei offenen Aufgaben kommt hinzu, seit wann sie offen sind, bei vorhandener Notiz der Zusatz mit Notiz. Eine Hauptaufgabe mit Unteraufgaben nennt außerdem, wie viele davon erledigt sind.

Auf einer Zeile schaltet die Eingabetaste zwischen offen und erledigt um, die Entf-Taste löscht die Aufgabe samt ihren Unteraufgaben, und die Kontextmenü-Taste bietet Als erledigt markieren beziehungsweise Wieder offen setzen, Titel bearbeiten, Notiz anzeigen, Notiz bearbeiten, Unteraufgabe hinzufügen und Löschen an. Eine Notiz zu leeren bedeutet, sie zu entfernen.

11.3. Zusammenspiel mit der KI

Die KI kann Aufgaben anlegen, auch gleich mit mehreren Unteraufgaben, zum Beispiel eine Einkaufsliste aus einem Rezept. Sie ist ausdrücklich angewiesen, Aufgaben, die sich beiläufig aus dem Gespräch ergeben, erst vorzuschlagen und nur nach deiner Bestätigung anzulegen. Bei einem direkten Auftrag legt sie sofort an. Auch das setzt den Schalter KI darf Aufgaben und Termine anlegen voraus.

12. Der Bereich Ollama Modelle

Hier verwaltest du die KI-Modelle auf deinem Rechner. Du erreichst den Bereich mit Strg plus Umschalt plus F10. Er ist immer bedienbar und an keine Freischaltung gebunden. Voraussetzung ist, dass Ollama läuft.

12.1. Ein Modell nachinstallieren

Ganz oben steht die Schaltfläche Modell-Installation öffnen. Im Fenster gibst du den Namen eines Modells aus der Ollama-Bibliothek ein, zum Beispiel qwen2 punkt 5 Doppelpunkt 7b, und drückst die Eingabetaste. Während des Downloads sagt das Programm jeweils in 10 Prozent Schritten an, wie viel bereits geladen ist. Modellnamen sind immer kleingeschrieben und enthalten einen Doppelpunkt; gibst du versehentlich einen Anzeigenamen ein, weist dich das Programm darauf hin. Daneben liegt ein Link, der die Ollama-Modellbibliothek im Browser öffnet. Wichtig ist hier, dass die Namen exakt so eingegeben werden, wie sie in der Ollama Bibliothek stehen. Zum Beispiel gema4 Doppelpunkt 26b. Ist mein aktueller Favorit! Ca. 18GB groß, und erfordert mindestens eine 16 GB VRAM RTX Grafikkarte.

12.2. Die Liste der installierten Modelle

Darunter folgt für jedes Modell ein eigener Eintrag. Beim Anspringen hörst du einen einzigen zusammenhängenden Satz mit allen Angaben: Name, Art, Größe, Kontextfenster, Kategorie, Status, Beschreibung und empfohlener Einsatz. Das ist Absicht, damit du nicht acht Einzelzeilen durchtabben musst.

Bei der Art wird zwischen Rohmodell und Modelfile-Variante unterschieden. Eine Variante nennt zusätzlich ihr Basismodell und den Hinweis, dass sie sich die Modelldaten mit dem Basismodell teilt und deshalb keinen zusätzlichen Speicher braucht.

In dieser Liste gibt es zwei Bewegungsarten. Die TAB-Taste geht durch alles hindurch, also Modell, dessen Schaltflächen, nächstes Modell. Die Pfeiltasten hoch und runter springen dagegen von Modell zu Modell und überspringen die Schaltflächen. Das ist der schnelle Weg, wenn du nur die Liste durchgehen willst. Pfeil hoch auf dem ersten Modell bringt dich zurück zur Überschrift Installierte Modelle, Pfeil runter auf dieser Überschrift wieder hinein.

12.3. Was du mit einem Modell tun kannst

Beschreibung bearbeiten: Bei jedem Modell vorhanden. Hier trägst du Kategorie, Beschreibung und empfohlenen Einsatz ein. Diese Angaben speichert das Programm für dich, sie verändern das Modell nicht.

Modelfile anlegen: Nur bei Rohmodellen. Damit machst du aus einem Rohmodell eine eigene, dauerhaft gespeicherte Variante mit festem Systemprompt und eigenen Werten. Mehr dazu im nächsten Abschnitt.

Systemprompt bearbeiten: Nur bei Modelfile Varianten. Ändert die feste Rolle und Arbeitsanweisung der Variante. Alle anderen Werte bleiben dabei erhalten.

Parameter bearbeiten: Nur bei Varianten. Ändert die technischen Werte der Variante, zum Beispiel das Kontextfenster.

Modell deinstallieren: Bei jedem Modell vorhanden. Es kommt eine Sicherheitsabfrage, und der Fokus steht dabei auf der Bestätigungs-Schaltfläche, die die vollständige Frage in ihrer Ansage trägt.

Achtung beim Deinstallieren eines Rohmodells: Alle davon abgeleiteten Modelfile Varianten werden mitgelöscht. Die Sicherheitsabfrage zählt sie dir vorher einzeln auf, und die Schaltfläche heißt dann sinngemäß Modell und zwei Varianten löschen.

12.4. Eigene Varianten anlegen

Eine Variante ist praktisch, wenn du dasselbe Modell für verschiedene Zwecke unterschiedlich einstellen möchtest, zum Beispiel eine kreative und eine sachliche Fassung. Die Variante startet dann jeden Chat automatisch mit deiner Rolle und deinen Werten, ohne dass du sie jedes Mal neu angeben musst. Das Rohmodell selbst bleibt unverändert, und die Variante braucht kaum zusätzlichen Speicherplatz.

Der Name darf Buchstaben, Ziffern, Punkt, Bindestrich und Unterstrich enthalten, wahlweise gefolgt von einem Doppelpunkt mit einer Versionsbezeichnung, also zum Beispiel gemma bindestrich rezepte, oder gemma bindestrich rezepte doppel punkt v1. Der Systemprompt ist die feste Rolle und Arbeitsanweisung, die jedem Chat mit dieser Variante vorangestellt wird.

Darunter gibt es zwölf technische Felder in vier Gruppen: Zufall und Vielfalt, Wiederholungen vermeiden, Mirostat sowie Länge, Kontext und

Reproduzierbarkeit. Jedes Feld erklärt in seiner Beschriftung, was es bewirkt, und nennt den Standardwert von Ollama. Felder, die du leer lässt, werden nicht geschrieben; dort gilt dann der Standard.

Kommazahlen kannst du mit Punkt oder Komma eingeben.

Ein Feld ist besonders wichtig: Kontextfenster in Tokens. Es bestimmt, wie viel Gesprächsverlauf das Modell gleichzeitig im Blick behält. Große Werte brauchen deutlich mehr Grafikspeicher. Um so mehr Grafikspeicher du hast, um so höher kann hier dein Wert ausfallen, auch hierbei kann dir eine KI gut helfen einen passenden Wert zu finden. Trägst du hier nichts ein, verwendet Ollama nur 2048 Tokens, Da dies der voreingestellte Standard wert von Ollama ist, es sei denn du hast diesen angepasst.

Die ganzen Einstellungen eines Modellfiles hören sich erst einmal sehr Technisch an, aber lass dir sonst von einer KI helfen welche Einstellungen du am besten für welche Zwecke verwenden kannst. Bei den jeweiligen Feldern wird dir aber auch von JAWS angesagt was ein guter Wert wäre. Ich habe versucht es so gut Verständlich wie Möglich zu gestalten Das Erstellen eines Modellfile kann je nach eingesetzter Verwendung wirklich gute Ergebnisse liefern.

13. Der Bereich KI-Verhalten

Hier siehst und änderst du die Anweisungen, mit denen das Programm die KI steuert. Du erreichst den Bereich mit Strg plus Umschalt plus F11. Er ist immer bedienbar.

13.1. Bitte zuerst lesen

Diese Anweisungen bestimmen, wie sich die KI grundsätzlich verhält. Ändere sie nur, wenn du dir sicher bist. Eine unklare, widersprüchliche oder zu knappe Anweisung kann dazu führen, dass die Modelle unerwartet reagieren, etwa die Sprache wechseln oder Anweisungen übergehen. Die gute Nachricht: Zu fast jedem Feld gibt es eine Schaltfläche Auf Standard zurücksetzen, mit der du jederzeit den Ursprungszustand wiederherstellst. Kaputtgehen kann dabei nichts.

Wichtig zu wissen: Nichts wird beim Tippen gespeichert. Erst die Speichern-Schaltfläche unter dem jeweiligen Feld schreibt deine

Änderung. Jedes Feld zeigt beim Öffnen den gerade wirksamen Text an, also entweder deinen gespeicherten oder den mitgelieferten Standard.

13.2. Die einzelnen Anweisungen

Alle Blöcke sind gleich aufgebaut: eine Überschrift, ein Erklärtext, das Eingabefeld und darunter die Schaltflächen. Das sind die Blöcke der Reihe nach:

System-Prompt für Chats: Die Grundanweisung für normale Chats und Projekt-Chats. Datum und Uhrzeit werden automatisch ergänzt und gehören nicht in dieses Feld.

Prompt für Bildschirmfoto: Wird mitgeschickt, wenn du mit Strg plus Umschalt plus B ein Bildschirmfoto aufnimmst, damit die KI den Bildschirm beschreibt.

Hinweis für Datei-Werkzeuge in Projekten: Wird in Projekt-Chats vorangestellt, in denen die KI Dateien anlegen und ändern darf, und weist sie an, die Werkzeuge wirklich zu benutzen.

Prompt für die Chat-Benennung: Steuert, wie der Titel eines neuen Chats automatisch entsteht.

Such-Entscheidungsregel der Websuche: Legt fest, wann das Modell bei eingeschalteter Websuche selbstständig sucht.

Anweisung für die Hörspiel-Szenenanalyse: Legt fest, wie das Modell Geräusche und Musikstimmung aus deinem Text heraussucht. Vorsicht: Das Modell muss dabei ein wörtliches Zitat liefern, mit dem das Programm die Sekunde in der Zeitbasis findet. Änderst du diese Regel unbedacht, kann die Verankerung fehlschlagen.

Anweisung für die Musikmomente: Legt fest, wie das Modell die Stellen für kurze Spannungsmusik findet. Die wichtigste Stellschraube ist die Sparsamkeit: Verlangst du zu viele Momente, schwillt ständig die Musik an, und das ermüdet beim Hören.

Begriffsliste für Übersetzungen: Feste Vorgaben, an die sich die KI beim Übersetzen und Schreiben immer halten soll, denn es kann mal Vorkommen das das eingesetzte KI Modell einen Text von Englisch ins Deutsche nicht richtig Übersetzt, wenn das Modell sich nicht sicher ist bleibt zum Beispiel ein Bestimmtes Wort einfach Englisch, und in dieser Liste kannst du das Problem Steuern. Eine Regel pro Zeile. Links der

Ausgangsbegriff, rechts die verbindliche Entsprechung, zum Beispiel Elf gleich Elb.

Car gleich Auto. Für Eigennamen, die unverändert bleiben sollen, schreibst du zum Beispiel Aeglos bleibt Aeglos. So heisst mein Schwert in meinen Geschichten, und es ist ein Wort was das Modell nicht kennt, und Übersetzt es in irgendwas passendes, und so kannst du dies unterdrücken. Statt eines Standards gibt es hier die Schaltfläche Begriffsliste leeren.

Danach folgen sechs Listen mit Auslöser-Wörtern für die automatische Modellwahl, nämlich für Coding, kreative Texte, Bildbeschreibung, Zusammenfassungen, Übersetzungen und OCR. Bei der ersten Nachricht eines Chats im Modus Automatisch prüft das Programm, welche dieser Wörter in deiner Frage vorkommen, und legt danach das passende Standardmodell fest. Schreibe ein Wort pro Zeile und klein. Verglichen wird, ob der Text irgendwo in deiner Frage vorkommt, auch als Teil eines längeren Wortes. Kommen Wörter aus mehreren Listen vor, gewinnt die Liste mit den meisten Treffern. Die Kategorie Allgemeine Chats hat bewusst keine Auslöser, sie greift immer dann, wenn nichts anderes eindeutig passt.

14. Der Bereich Einstellungen

Hier stellst du das Programm auf dich ein. Du erreichst den Bereich mit Strg plus Umschalt plus F12. Der Fokus landet dabei direkt auf der Schaltfläche Lizenz anzeigen und verwalten ganz oben; von dort führt die TAB-Taste weiter zur Update-Schaltfläche und danach zu den eigentlichen Einstellungen. Alle Änderungen wirken sofort und werden sofort gespeichert, es gibt keine Speichern-Schaltfläche.

Bei den Zahlenfeldern gilt eine gemeinsame Regel: Sie nehmen nur Ziffern an, du kannst den Wert mit den Pfeiltasten hoch und runter verändern, und beim Verlassen wird ein Wert außerhalb des erlaubten Bereichs zurechtgerückt und die Korrektur angesagt. Schieberegler gibt es im ganzen Programm bewusst keine, weil ein Screenreader deren Stellung nur als Prozentwert vorliest.

Die folgenden Abschnitte entsprechen genau den Abschnitten im Programm, in der Reihenfolge, in der du sie beim Tabben erreichst. In Klammern steht jeweils der Wert ab Werk.

14.1. Konto und Lizenz

Ganz oben im Einstellungs-Bereich, noch vor allen Einstellungen, stehen zwei Schaltflächen nebeneinander. Sie betreffen dein Konto und die Programmfassung.

Lizenz anzeigen und verwalten

Diese Schaltfläche öffnet ein Fenster mit dem Namen Lizenz-Status und Geräteverwaltung. Unter der Überschrift Deine Lizenz siehst du den aktuellen Stand deiner Lizenz. Darin gibt es drei Schaltflächen:

Schlüssel erneut prüfen: Fragt beim Server nach und aktualisiert deinen Lizenz-Status. Das ist der richtige Weg, wenn du deine Lizenz gerade verlängert oder erweitert hast und das Programm noch den alten Stand zeigt.

Gerät abmelden: Meldet genau diesen Rechner von deiner Lizenz ab und gibt den belegten Aktivierungsplatz wieder frei. Das brauchst du, wenn du auf einen neuen Rechner umziehst. Weil das nicht versehentlich passieren soll, musst du zuvor ein Kontrollfeld anhaken, das ausdrücklich bestätigt, dass du dieses Gerät abmelden und den Platz freigeben möchtest.

Kundenportal öffnen: Öffnet das Kundenportal im Browser. Dort siehst du deine Schlüssel und Geräte, kannst Support anfordern und ein Anliegen melden.

Mit Schließen kommst du zurück ins Programm.

Auf Updates prüfen

Diese Schaltfläche prüft online, ob für deine Lizenz eine neue Programmfassung bereitsteht. Wichtig zu wissen: Das Programm sucht nicht heimlich im Hintergrund nach Updates. Es prüft nur, wenn du diese Schaltfläche drückst. Über neue Updates wirst du stattdessen über die Broadcast Benachrichtigungen informiert, siehe Kapitel 16.

Steht ein Update bereit, öffnet sich ein Fenster mit dem Namen Programm-Updates. Darin steht unter der Zeile Das ist neu in dieser Version eine Zusammenfassung der Änderungen. Darunter findest du diese Schaltflächen:

Update herunterladen und installieren: Lädt die neue Fassung herunter, prüft die Datei auf Unversehrtheit und installiert sie. Das Prüfen auf Unversehrtheit, heisst nur das hier ein Manipulationsschutz eingebaut ist, so das bei einem Server Einbruch diese Datei nicht geändert werden kann, und somit bei euch zum Beispiel Malware Installiert wird. Der Server ist in der Regel gut gesichert, aber einen hundert Prozentigen Schutz gibt es nie. Das Programm startet danach automatisch neu Sicherheits Updates werden gesondert geliefert, und nicht über das Programm. Da Sicherheits-Updates alle erhalten, egal ob ihr Update Paket abgelaufen ist oder nicht. Sollte euer Update Paket abgelaufen sein, und ihr seid somit für ein Update nicht Berechtigt, werdet ihr darüber informiert. Während des Downloads kannst du den Vorgang abbrechen; wird die Prüfung der Datei nicht bestanden, wird die Datei gelöscht und nichts installiert.

Erneut prüfen: Fragt noch einmal beim Server nach.

Vollständige Patch Notes öffnen: Öffnet die ausführliche Liste aller Änderungen dieser Fassung im Browser.

Kundenportal öffnen: Öffnet das Kundenportal. Dort kannst du deine Lizenz einsehen, Downloads abrufen und den Update-Zeitraum verlängern lassen. Hier kannst du auch einen Downloadlink erzeugen lassen, dieser Temporärer Link ist dann für 24 Stunden gültig. Im Kundenportal kannst du aber jederzeit einen neuen erhalten. Denn die allgemeinen Updates sind nicht im öffentlichen Webroot platziert sondern in einer gesicherten Umgebung, so das keine Bots und unberechtigte Personen an diese dran kommen.

Die erstmalige Aktivierung

Beim allerersten Start fragt dich das Programm nach deinem Lizenzschlüssel. Das Fenster heißt Programm aktivieren und hat zwei Felder. In das erste trägst du deinen Lizenzschlüssel ein; er beginnt mit THAL, gefolgt von vier Gruppen mit je fünf Zeichen. Die Bindestriche kannst du weglassen. Kannst sie aber auch verwenden, es geht also mit und ohne. In das zweite Feld trägst du einen Gerätenamen ein, unter dem dieser Rechner in deiner Geräteverwaltung im Kundenportal erscheint, zum Beispiel Arbeits-PC. Hier wird der Voreingestellte Wert deines Windows PC genommen, je nachdem was du im Windows Eingestellt hast

bei Computernamen. Aber diesen Wert kannst du anpassen wenn dir ein anderer Name lieber ist.

Für diese erstmalige Aktivierung brauchst du eine Internetverbindung. Danach funktioniert das Programm auch ohne Internet, denn die Lizenz wird verschlüsselt auf deinem Rechner hinterlegt und dort geprüft. Ohne einen gültigen Schlüssel wirst du nicht in der Lage sein, das Programm zu starten und auszuführen; es wird sich stattdessen direkt wieder schließen. Bei zu vielen fehlerhaften Eingaben wird zudem die Nutzung des Programms komplett gesperrt. Diese Sperre lässt sich nicht umgehen, sodass in einem solchen Fall nur der Kundendienst weiterhelfen kann. Mit der „Sperre“ ist gemeint, dass du das Programm überhaupt nicht mehr starten kannst – auch das Aktivierungsfenster bleibt dann unzugänglich. Dein PC selbst bleibt davon natürlich weiterhin voll funktionsfähig..

14.2. Ansprechnamen

Ansprechnamen: Dein Name, mit dem das Programm und die KI dich ansprechen. Bleibt das Feld leer, gibt es eine allgemeine Begrüßung ohne Namen. (Vorgabe: leer)

14.3. Verbindung zu Ollama

Ollama URL: Die Adresse, unter der das Programm Ollama erreicht. Nur ändern, wenn Ollama auf einem anderen Rechner oder Anschluss läuft. (Vorgabe: <http://127.0.0.1:11434>). Dies ist der voreingestellte Standardwert von Ollama, ändere diesen Wert daher nur wenn du weißt was du da tust. In der Regel muss es nicht geändert werden.

14.4. Standardmodelle

Für jede Art von Aufgabe kannst du ein eigenes Modell festlegen. Bleibt ein Feld leer, wird das gerade ausgewählte Hauptmodell verwendet. Alle Felder sind Auswahlfelder und ab Werk leer.

Allgemeine Chats: Das Modell für normale Unterhaltungen. Dieses Modell übernimmt auch die Pflege der Erinnerungen, wenn kein Modell für Zusammenfassungen gewählt ist.

Coding und Programmierung: Für Programmier-Fragen und Code-Aufgaben.

Kreative Texte: Für Geschichten und kreatives Schreiben, auch im Bereich Kreatives.

Bildbeschreibung: Für das Beschreiben von Bildern. Dafür ist ein bildfähiges Modell nötig.

Zusammenfassungen: Für das Zusammenfassen von Texten und für die Pflege der Erinnerungen.

Übersetzungen: Für Übersetzungen, auch bei der Transkription.

Blockgröße für Transkript-Übersetzung in Zeichen: Wie große Stücke beim Übersetzen eines Transkripts am Stück bearbeitet werden. Das bestimmt zugleich die Absatzlänge im fertigen Dokument. Wirkt nur, wenn weder Zeitstempel noch Sprechertrennung eingeschaltet sind. (Vorgabe: 10000, erlaubt 1000 bis 30000)

OCR: Für die Texterkennung in Bildern.

Chat-Benennung: Das Modell, das im Hintergrund einen kurzen Titel für einen neuen Chat vergibt. Ein kleines, schnelles Modell genügt hier völlig.

Hörspiel-Szenenanalyse: Das Modell, das deinen Hörspieltext liest und daraus Geräusche und Musikstimmung heraussucht. Bleibt es leer, wird das Modell für allgemeine Chats verwendet.

Einbettungen: Das Modell, mit dem das Programm die passenden Erinnerungen zu deiner Frage findet. Bleibt das Feld leer, sucht das Programm selbst das erste installierte Modell, dessen Name das Wort embed enthält. Wichtig: Modelle ohne dieses Wort im Namen musst du hier ausdrücklich auswählen, sonst werden sie nicht gefunden.

14.5. Chat-Verhalten

Denkprozess der KI im Chat anzeigen: Zeigt die internen Denk-Abschnitte mancher Modelle als eigenen Bereich im Chat. Ausgeschaltet bedeutet: Nur die eigentliche Antwort erscheint. Die Qualität bleibt gleich. (Vorgabe: aus)

Nachdenken der Modelle nur in Projekten erlauben: Schaltet das Nachdenken außerhalb von Projekt-Chats ab, auch bei Hilfsaufgaben wie der Chat-Benennung. Das beschleunigt die Antworten spürbar. (Vorgabe: ein). Nicht jedes Modell in der Ollama-Bibliothek ist ein Reasoning-Modell. Solche Modelle eignen sich hingegen sehr gut für

Coding-Projekte, da Modelle, die über „Nachdenkprozesse“ verfügen, oft besseren Code liefern können. Zudem denken diese Modelle häufig auf Englisch nach, was man leider nicht beeinflussen kann.

14.6. Kontext-Warner

Der Kontext ist der Gesprächsspeicher des Modells. Ist er voll, gehen die ältesten Teile des Gesprächs verloren. Diese vier Einstellungen bestimmen, wann das Programm dich warnt.

Maximale Kontext-Tokens des aktuellen Modells: Wie viel dein Modell fassen kann. Dieser Wert dient als Bezugsgröße für die Warnungen und wird zugleich an Ollama weitergegeben, sofern das Modell keinen eigenen Wert mitbringt. (Vorgabe: 8192, erlaubt 2048 bis 131072). Auf Blindsoft.de findest du im Bereich der Dokumentationen verschiedene passende Werte, je nachdem, wie viel Grafikspeicher (VRAM) deine Hardware besitzt. Grundsätzlich gilt: Je mehr Grafikkartenspeicher du verfügst, desto höher kann der CTX-Wert eingestellt werden.

Hier sind einige Richtwerte für die Konfiguration nach Speicherkapazität:

Empfohlene Werte basierend auf dem VRAM

- 4 GB VRAM: ca. 2048 bis 4096 (2K bis 4K)
- 8 GB VRAM: ca. 8192 (8K)
- 12 GB VRAM: ca. 16384 (16K)
- 16 GB VRAM: ca. 32768 (32K)
- 24 GB VRAM: ab 49152 (48K) bis 65536 (64K) und höher.

Der hier eingestellte Wert hat stets Vorrang, ungeachtet der Einstellungen, die du in Ollama selbst vorgenommen hast. Solltest du jedoch ein Modelfile verwenden, in dem bereits ein CTX-Wert explizit hinterlegt ist, so hat dieser Wert Vorrang vor

der hier getroffenen Einstellung. Der hier konfigurierbare Wert dient also primär für Rohmodelle oder Modelfiles, die keinen eigenen CTX-Wert spezifiziert haben.

Schwellwert niedrig in Prozent: Ab dieser Auslastung kommt der erste Hinweis, dass du bald einen neuen Chat beginnen solltest. (Vorgabe: 70)

Schwellwert mittel in Prozent: Ab hier kommt die zweite, deutlichere Warnung. (Vorgabe: 85)

Schwellwert hoch in Prozent: Ab hier wird gewarnt, dass wichtige Inhalte verloren gehen könnten. (Vorgabe: 95)

14.7. Erinnerungen

Erinnerungen speichern: Die KI bekommt bei jeder Anfrage die dauerhaft gemerkten und die zur Frage passenden Erinnerungen mit. Der Schalter steuert auch, ob sie sich selbst etwas merken darf. (Vorgabe: ein)

Erinnerungen beim Beenden automatisch aktualisieren: Beim Schließen des Programms werden die Erinnerungen aus den Chats gepflegt. Beginn und Ende werden angesagt, und das Programm schließt sich erst danach. Wirkt nur, wenn Erinnerungen speichern eingeschaltet ist. (Vorgabe: ein)

Beispielprompt öffnen: Zeigt dir einen empfohlenen Text, mit dem du deine Erinnerungen aus einem anderen KI-Dienst herausholen kannst.

Text aus dem Export einfügen: Hier fügst du die aus einem anderen Programm exportierten Erinnerungen ein, eine Zeile je Erinnerung.

Importieren: Übernimmt die eingefügten Erinnerungen. Doppelte werden übersprungen.

Erinnerungen anzeigen: Öffnet die Verwaltung deiner gespeicherten Erinnerungen.

Erinnerungen exportieren: Gibt alle Erinnerungen im Import-Format aus, wahlweise in die Zwischenablage oder als Textdatei.

Erinnerungen aus Chats aktualisieren: Wertet die neuen Nachrichten aller Chats aus und pflegt die Erinnerungen. Das ist derselbe Vorgang, der sonst beim Beenden läuft, und kann einen Moment dauern.

Das Fenster hinter der Schaltfläche Erinnerungen anzeigen

Dieses Fenster trägt den Namen Erinnerungen. Darin steht jede gespeicherte Erinnerung in einer eigenen Zeile, und zu jeder Zeile gehören zwei Bedienelemente.

Das erste ist ein Kontrollfeld. Damit legst du fest, ob diese Erinnerung immer aktiv ist. Eine immer aktive Erinnerung wird jeder einzelnen Anfrage beigelegt, ganz gleich worum es geht. Sie eignet sich für dauerhafte Vorgaben über dich oder deine Arbeitsweise. Ist das Kontrollfeld nicht angehakt, ist die Erinnerung situativ: Sie wird nur dann mitgeschickt, wenn sie inhaltlich zu deiner Frage passt. Für die meisten Erinnerungen ist das die richtige Einstellung, denn so bleibt der Gesprächsspeicher frei für das Wesentliche.

Das zweite ist eine Schaltfläche zum Löschen dieser einen Erinnerung. Nach dem Löschen springt der Fokus gezielt auf die nächste Erinnerung, du kannst also zügig aufräumen.

Über der Liste gibt es ein Eingabefeld für eine neue Erinnerung und daneben die Schaltfläche Erinnerung hinzufügen. So trägst du etwas von Hand nach, ohne es der KI im Chat sagen zu müssen. Ganz oben steht außerdem die Schaltfläche Alle Erinnerungen löschen, die vorher sicherheitshalber nachfragt.

14.8. Websuche

Brave API-Key: Dein Zugangsschlüssel für die Suchmaschine Brave. Die Eingabe wird nicht angezeigt, und der Schlüssel wird verschlüsselt gespeichert. (Vorgabe: leer)

Websuche aktivieren: Erlaubt der KI, im Internet zu suchen. (Vorgabe: ein)

Vor Websuche nachfragen: Vor jeder Suche erscheint eine Rückfrage mit den Suchbegriffen und deinem Anfragestand des Monats. (Vorgabe: ein).

Brave bietet eine Suchmaschine für KIs an. Du kannst dort einen kostenlosen API-Schlüssel erhalten, mit dem dir jeden Monat ein Guthaben von 5 US-Dollar zur Verfügung steht. Damit sind monatlich etwa 1000 Anfragen an Brave möglich. Um dich vor unerwarteten Kosten zu schützen, ist hier standardmäßig ein entsprechend limitierter Wert voreingestellt. Wie du einen eigenen Brave API-Schlüssel erstellst, erfährst du in den anderen Dokumentationen, die dem Programm beiliegen, oder auf blindsoft.de. **-Anfrage-Kontingent:** Wie viele Anfragen dein Brave-Vertrag pro Monat erlaubt. Dient nur als Bezugswert für die Anzeige. (Vorgabe: 1000)

Anzahl der Web-Treffer pro Suche: Wie viele Treffer der KI übergeben werden. Mehr Treffer geben mehr Stoff, verbrauchen aber mehr Platz im Gespräch. (Vorgabe: 8, erlaubt 1 bis 20)

Maximale Anzahl Suchbegriffe pro Websuche: Obergrenze für die Suchbegriffe, die die KI aus deiner Frage bildet. Eine kurze Suchanfrage liefert meist die besten Ergebnisse. (Vorgabe: 6, erlaubt 2 bis 12)

Maximale Anzahl Websuchen pro Antwort: Wie oft die KI während einer einzigen Antwort selbstständig suchen darf. Das schützt dein Kontingent. (Vorgabe: 3, erlaubt 1 bis 10)

Unter diesen Feldern zeigt das Programm deinen aktuellen Anfragestand des Monats an.

Was beim Nachfragen genau passiert

Ist die Einstellung Vor Websuche nachfragen eingeschaltet, öffnet sich vor jeder einzelnen Suche ein kleines Fenster. Darin steht eine Frage, die zwei Dinge nennt: erstens die Suchbegriffe, die die KI aus deiner Frage gebildet hat, und zweitens deinen aktuellen Anfragestand des Monats, also wie viele deiner Anfragen bereits verbraucht sind. So siehst du vor jeder Suche, wonach gesucht werden soll und was es dich kostet.

Es gibt zwei Schaltflächen: Websuche starten und Websuche nicht starten. Lehnst du ab, läuft deine Anfrage trotzdem weiter, die KI antwortet dann eben ohne Internet. Nach dem Schließen des Fensters springt der Fokus gezielt zurück in das Eingabefeld des Chats.

Wichtig zu wissen: Die Nachfrage gilt für jede einzelne Suche. Sucht die KI während einer Antwort mehrmals, wirst du auch mehrmals gefragt.

Wie oft sie das überhaupt darf, begrenzt die Einstellung Maximale Anzahl Websuchen pro Antwort. Schaltest du die Nachfrage aus, sucht die KI ohne Rückfrage, und dann schützt dich nur noch diese Obergrenze vor unnötigem Verbrauch.

Diesen Schalter habe ich zur Sicherheit eingebaut, da der Prompt für die Websuche und das dahinterliegende System derzeit noch nicht so präzise funktioniert, wie ich es mir wünsche. Es kommt gelegentlich vor, dass ein Modell eine Websuche startet, die inhaltlich keinen Sinn ergibt oder schlichtweg unnötig ist. Um dein Kontingent zu schützen, dient dieser Sicherheitsschalter, bei dessen Aktivierung du erst um Erlaubnis gefragt wirst. Falls du diese Funktion also bewusst ausgeschaltet hast, stelle dich darauf ein, dass das Modell dennoch im Internet sucht, selbst wenn dies für deine eigentlich gestellte Frage gar nicht erforderlich gewesen wäre.

Bedenke zudem, dass nicht jedes KI-Modell in der Ollama-Bibliothek über sogenannte „Tools“ verfügt. Wenn ein Modell beispielsweise nicht über Funktionen wie „Web Search“ oder „Function Calling“ verfügt, kann es sein, dass die Websuche nicht einwandfrei funktioniert. Auf blindsoft.de findest du eine entsprechende Liste mit Ollama-Modellen, die sich gut verwenden lassen und besonders für Grafikkarten mit weniger Speicherkapazität geeignet sind.

14.9. Webseiten-Abruf

Webseiten-Abruf aktivieren: Schickst du im Chat eine Internetadresse mit, holt das Programm die Seite selbst und übergibt der KI den Haupttext, damit sie ihn zusammenfassen kann. Abgerufen wird nur die erste Adresse einer Nachricht. (Vorgabe: ein)

14.10. Sprachausgabe mit der Windows-Vorlesestimme

Hier geht es um die eingebauten Windows-Stimmen, nicht um die KI-Stimme aus Kapitel 7.

Vorlesestimme: Die Windows-Stimme, mit der Antworten vorgelesen werden. Die Liste zeigt die bei dir installierten Stimmen. (Vorgabe: Katja, sofern vorhanden)

Probe anhören: Spielt einen kurzen Beispielsatz mit der gewählten Stimme.

Antworten automatisch vorlesen: Jede fertige Antwort wird selbsttätig vorgelesen. Ist das eingeschaltet, hält sich dein Screenreader dabei zurück, damit nicht zwei Stimmen gleichzeitig sprechen. (Vorgabe: aus).

Diese Funktion hat keine Auswirkungen auf die „Push-to-Talk“-Funktion. Wenn diese ausgeschaltet ist und du „Push-to-Talk“ mittels STRG + UMSCHALT + LEERTASTE nutzt, um einen Text zu diktieren, wird die Antwort der KI stets mit der voreingestellten Stimme abgespielt. JAWS bleibt in diesem Fall wie gesagt stumm. Die Funktion dient dazu, dass du weiterhin mit der KI interagieren kannst, selbst wenn das Programm minimiert ist und nicht im Vordergrund steht.

Vorlese-Geschwindigkeit: Wie schnell die Windows-Stimme spricht. Null ist normal. (Vorgabe: 0, erlaubt minus 10 bis plus 10)

14.11. Audio Transkribierung

Hugging-Face-Token: Der Zugangsschlüssel, der ausschließlich für die Sprechertrennung gebraucht wird. Ohne ihn entstehen Transkript und Zeitstempel trotzdem, nur die Sprecher werden nicht getrennt. Die Eingabe wird nicht angezeigt und verschlüsselt gespeichert. Wie du einen API-Schlüssel auf Hugging Face erstellen kannst, erfährst du in der Dokumentation 2, die diesem Programm beiliegt, oder auf <https://blindsoft.d> e. (Vorgabe: leer)

14.12. Werkzeuge der KI

Maximale Anzahl KI-Werkzeug-Aufrufe pro Antwort: Wie viele Werkzeuge die KI in einer einzigen Antwort verwenden darf, etwa beim Anlegen mehrerer Dateien. Bei größeren Projekten ist ein höherer Wert sinnvoll. (Vorgabe: 25, erlaubt 1 bis 100)

Erlaubte Kommandozeilen-Befehle: Die Liste der Befehle, die die KI in Programmier-Projekten ausführen darf, ein Befehl je Zeile. Vor jeder Ausführung wirst du ohnehin um Bestätigung gebeten; Befehle außerhalb dieser Liste bekommen zusätzlich einen Warnhinweis. Ab Werk ist eine Liste gängiger Entwickler-Befehle eingetragen. Leerst du das Feld, wird beim nächsten Start die Standardliste wiederhergestellt.

14.13. Anzeige und Ansagen

Schriftgröße: Die Schriftgröße im ganzen Programm. Die Änderung wirkt sofort. (Vorgabe: 21, erlaubt 14 bis 32)

Intervall der Fortschritts-Ansagen in Sekunden: In welchem Abstand das Programm während eines laufenden Vorgangs eine Zwischenansage macht, damit du weißt, dass weitergearbeitet wird. (Vorgabe: 10, erlaubt 5 bis 30)

14.14. Programm-Start

Begrüßungsbildschirm anzeigen: Zeigt beim Start einen Begrüßungstext. (Vorgabe: ein)

Programm minimiert in Hintergrund starten: Das Hauptfenster bleibt beim Start minimiert in der Taskleiste. (Vorgabe: aus)

Mit Windows minimiert im Hintergrund starten: Trägt das Programm in den Windows-Autostart ein. Nach der Anmeldung startet es selbsttätig und minimiert. Mit Strg plus Alt plus T holst du es nach vorn. (Vorgabe: aus)

Wartezeit auf Ollama beim Windows-Autostart in Sekunden: Beim Start über den Autostart wartet das Programm so lange, bevor es Ollama prüft, denn Ollama braucht nach der Anmeldung meist etwas länger. (Vorgabe: 30, erlaubt 0 bis 300)

14.15. Assistenz

Assistenz-Fenster beim Start anzeigen: Zeigt beim Start ein Fenster mit deinen offenen Aufgaben und Terminen. Näheres in Kapitel 15. (Vorgabe: aus)

Anzahl der Aufgaben im Assistenz-Fenster: Wie viele offene Aufgaben das Fenster höchstens zeigt. (Vorgabe: 3, erlaubt 1 bis 20)

Anzahl der Termine im Assistenz-Fenster: Wie viele der nächsten Termine das Fenster höchstens zeigt. Überfällige Termine werden immer alle gezeigt. (Vorgabe: 3, erlaubt 1 bis 20)

KI darf Aufgaben und Termine anlegen: Die KI kann auf deine Bitte Aufgaben und Termine anlegen und von sich aus welche vorschlagen. Eingetragen wird erst nach deiner Bestätigung. (Vorgabe: ein)

14.16. Termin-Erinnerung

Termin-Erinnerungen aktiviert: Das Programm meldet sich bei Fälligkeit eines Termins mit Fenster und Ton. Ist es minimiert, holt es sich in den Vordergrund. (Vorgabe: ein)

Eigener Ton als WAV-Datei: Pfad zu einer eigenen Klangdatei für die Erinnerung. Bleibt das Feld leer oder fehlt die Datei, kommt ein schlichter Standardton. (Vorgabe: leer)

Durchsuchen: Öffnet die Dateiauswahl für diese WAV-Datei.

Ton testen: Spielt den eingestellten oder den Standardton ab.

14.17. Backup

Automatische Backups: Legt eine Sicherungskopie an, bevor die KI in einem Projekt eine Datei ändert oder löscht. (Vorgabe: aus)

Backup-Speicherpfad: Der Ordner für die Sicherungskopien. Dieses Feld erscheint nur, wenn die automatischen Backups eingeschaltet sind. Bei leerem Feld wird ein Standardordner verwendet. (Vorgabe: leer)

14.18. Speicherpfade

Hinter jedem dieser Felder steht eine Schaltfläche Durchsuchen, die die Ordnerauswahl öffnet. Alle Felder sind ab Werk leer; das Programm verwendet dann jeweils einen sinnvollen Standardordner.

Speicherpfad für Projekte: Der Ordner, unter dem alle Projekt-Ordner angelegt werden. Ohne diesen Pfad kannst du keine Projekte anlegen.

Speicherpfad für Übergaben und Transkripte: Ablage für Übergabe-Protokolle und Transkripte aus Chats ohne Projekt. Bei Projekt-Chats landen sie stattdessen im Projekt.

Speicherpfad für Geschichten: Ordner für die Geschichten aus dem Bereich Kreatives, jede in einem eigenen Unterordner.

Speicherpfad für generierten Code: Ablage für Archive mit von der KI erzeugtem Programmcode.

Speicherpfad für KI-Stimmen und geklonte Stimmen: Ablage für erzeugte Stimmen. Die Referenzstimmen liegen im Unterordner Referenzstimmen darunter.

Speicherpfad für eigene Klänge des Hörspiels: Ordner mit deinen eigenen Klangdateien, die du im Text mit einer Markierung aufrufst. Auch die Werkstatt speichert hierhin.

Speicherpfad für eigene Musik des Hörspiels: Dasselbe für eigene Musikstücke.

Speicherpfad für fertige Hörspiele: Wohin die Hörspiel-Projektordner gelegt werden.

Hörspiele aufräumen nach Tagen: Nach wie vielen Tagen die großen Zwischendateien alter Hörspiele entfernt werden. Das fertige Hörspiel und dein Originaltext bleiben erhalten, und die fünf neuesten Ordner werden nie angetastet. (Vorgabe: 30, erlaubt 1 bis 365)

Speicherpfad für Dokumente: Ablage für Word-Dokumente, die die KI in Chats ohne Projekt erzeugt.

14.19. Logging

Aufbewahrungszeit für Log-Dateien in Tagen: Wie lange die Protokolldateien aufgehoben werden. (Vorgabe: 30, erlaubt 7 bis 365)

Debug-Modus für Log-Einträge aktivieren: Schreibt zusätzliche, sehr ausführliche Einträge. Das ist bei der Fehlersuche hilfreich, sollte im Alltag aber ausgeschaltet bleiben. (Vorgabe: aus)

Jetzt bereinigen: Räumt die alten Protokolldateien sofort auf.

Log-Archiv öffnen: Öffnet den Ordner mit den Protokolldateien im Explorer. Genau diese Dateien schickst du dem Support, wenn etwas nicht funktioniert.

14.20. Modell-Updates

Hier geht es ausdrücklich um die Modelle der Audio-Kette, also Musik, Geräusche, Sprache und Transkription, nicht um Programm-Updates und nicht um deine Ollama-Modelle.

Auf Updates prüfen: Sucht neuere Fassungen dieser Modelle und lädt nur das herunter, wozu es wirklich etwas Neues gibt. Die alte Fassung wird danach entfernt.

Abbrechen: Bricht eine laufende Prüfung ab.

14.21. Anhang-Dateien aufräumen

Anhang-Dateien jetzt aufräumen: Entfernt alle angehängten Dateien, auf die keine Nachricht mehr verweist, also Reste gelöschter Chats und nie gesendete Anhänge. Das Programm räumt bei jedem Start ohnehin selbst auf, diese Schaltfläche ist nur für den Fall, dass du sofort Platz schaffen möchtest.

15. Das Assistenz-Fenster beim Start

Wenn du in den Einstellungen das Assistenz-Fenster eingeschaltet hast, erscheint es beim Programmstart und gibt dir einen kurzen Überblick über deinen Tag. Es kommt nach dem Begrüßungsbildschirm oder, wenn dieser ausgeschaltet ist, sofort.

Das Fenster begrüßt dich mit deinem Ansprechnamen und sagt dann in zwei Sätzen, wie viele unerledigte Aufgaben und wie viele anstehende sowie überfällige Termine du hast. Danach kommt der Hinweis, dass du mit der TAB-Taste die Einträge durchgehen und mit der Eingabetaste einen davon öffnen kannst.

Darunter stehen die beiden Abschnitte Aufgaben und Termine mit den einzelnen Einträgen. Wie viele angezeigt werden, legst du in den Einstellungen fest; überfällige Termine erscheinen immer alle. Gibt es mehr Einträge als angezeigt, steht darunter der Hinweis, dass du den Rest im jeweiligen Bereich findest.

Der Fokus liegt beim Öffnen auf der Schaltfläche Schließen. Das ist Absicht: So liest dein Screenreader den Inhalt genau einmal vor, ohne ihn beim Weitertabben zu wiederholen. Drückst du auf einem Eintrag die Eingabetaste, schließt sich das Fenster, das Programm wechselt in den

passenden Bereich und setzt den Fokus genau auf diesen Eintrag. Schließt du einfach nur, landest du im Chat-Eingabefeld.

Das Assistenz-Fenster ist die einzige Stelle, an der dich das Programm auf Termine hinweist, die schon vor dem Start überfällig waren. Die laufende Termin-Erinnerung deckt nur Fälligkeiten während des Betriebs ab.

16. Weitere Fenster

16.1. Übersicht der Tastenbefehle

Mit Strg plus Umschalt plus H öffnest du jederzeit die vollständige Liste aller Tastenbefehle. Sie wird direkt aus dem Programm erzeugt und ist damit immer aktuell.

16.2. Rückmeldung an die Entwicklung

Mit Strg plus Umschalt plus F öffnest du ein Fenster, mit dem du eine Rückmeldung senden kannst, etwa einen Fehlerbericht oder einen Wunsch.

16.3. Benachrichtigungen

Alle meine Programme verfügen über ein „Broadcast-Mitteilungssystem“. Das bedeutet, dass ich über mein WebUI eine Nachricht an ein ausgewähltes Programm senden kann. Wenn es beispielsweise ein Update gibt, enthält diese Mitteilung wichtige Informationen sowie einen Link zu den Patch Notes und alles Weitere, was dazu gehört. Sofern eine Internetverbindung besteht, erscheint diese Meldung direkt beim Starten des Programms.

Dieses System wird für Ankündigungen von Updates und Sicherheitsupdates genutzt, aber auch für Rabattaktionen und Gewinnspiele. Es kann vorkommen, dass ihr eine Zeitlizenz für ein Programm gewinnt oder sogar eine dauerhafte Lizenz, die ihr anschließend an jemanden verschenken könnt. Das System wird jedoch so sparsam wie möglich eingesetzt, um euch nicht mit unnötigen Mitteilungen zu belästigen.

Mit Strg plus Umschalt plus I kannst du jederzeit die Benachrichtigungen von Thaliruth Arda Intelligence erneut abrufen. Darüber wirst du unter anderem über neue Programmfassungen informiert.