

Anleitung um legendäre Hooks mit Claude + Higgsfield zu erstellen

1) Higgsfield Account erstellen

Hier Account erstellen: <https://wperfolg.de/higgsfield>

2) Higgsfield MCP mit Claude verknüpfen

Hier findest Du die MCP URL: <https://higgsfield.ai/mcp>

3) Videoszene aufnehmen

Wir setzen uns z.B. auf einen Stuhl und sprechen den Text der Szene ein. Das kannst Du auch mit einer ganz normalen Handykamera filmen. Du brauchst nichtmal zwingend ein Stativ. Sinnvoll wenn sich Deine Kleidung vom Hintergrund abhebt – Weißes Hemd mit weißem Hintergrund ist nicht optimal.

→ Ggf. Video kostenfrei zuschneiden mit z.B. CapCut: <https://wperfolg.de/capcut>

4) Wir machen einen Screenshot von der ersten Szene des Videos

Dieser Screenshot ist sehr wichtig. Ihn nutzen wir jetzt um uns in eine Gefahrensituation zu setzen.

5) Prompt 1: Bilder in Gefahrensituationen generieren lassen

Verwende Nano Banana Pro über den Higgsfield MCP. Auflösung [1K]. Seitenverhältnis [9:16 - Reels].

Lasse die Person im angehängten Bild genau gleich (Komposition). Entferne alle Elemente [außer XY].

Gib mir 4 Bilder, welche die Person jeweils in einer von 4 unterschiedlichen Gefahrensituation darstellen. [Eine Situation soll sein: Ein Löwe rennt auf die Person zu].

Das Bild wird später der Startframe für ein Video, das [X] Sekunden lang sein wird. Die Gefahr soll erst dann die Person erreichen, wenn die entsprechende Anzahl an Sekunden vorbei ist. Achte also darauf, beim Erstellen der Bilder.

6) Prompt 2: Videoprompt generieren lassen

Du bist ein Experte für Videoprompts für KI-Videomodelle in Higgsfield (z. B. Seedance, Kling, Veo). Erstelle mir einen perfekten, produktionsreifen Videoprompt in englischer Sprache.

SCHRITT 1 — Frage mich ZUERST alle folgenden Angaben ab (alle auf einmal, kompakt als Liste), bevor du irgendetwas anderes tust. Falls ich einzelne Angaben bereits in meiner Nachricht mitgeliefert habe, übernimm sie und frage nur die fehlenden ab:

- 1. Videomodell? (Standard-Vorschlag: Seedance 2.0 — bestätigen oder anderes nennen)*
- 2. Videolänge in Sekunden?*
- 3. Name/Referenz des Startbilds? (z. B. "startbild-fertig" — wird als exakter erster Frame verwendet)*
- 4. Name/Referenz des Gestik-Referenzvideos? (z. B. "video- v1")*
- 5. Was passiert im Hintergrund / welches Ereignis oder welche Gefahr nähert sich? (z. B. Lawine, Tornado, Hai ...)*
- 6. Wie soll der Endzustand bei der letzten Sekunde sein? (Standard: Die Gefahr ist bedrohlich nah, direkt hinter der Person, verschluckt sie aber NICHT — Person bleibt voll sichtbar)*
- 7. Kamera? (Standard: statisch, ein durchgehender Shot, keine Schnitte, keine Kamerabewegung)*
- 8. Redet die Person durchgehend? (Standard: ja, ohne Unterbrechung von Sekunde 0 bis zum Ende)*
- 9. Videoqualität und Videoformat? (Standard ist 720p und 9:16 (Reels))*

SCHRITT 2 — Erstelle danach den finalen Videoprompt auf Englisch nach genau diesem Aufbau:

a) *STARTFRAME*: Das angegebene Startbild wird als exakter erster Frame referenziert. Gesamtlänge, ein durchgehender Shot, Kameraverhalten laut Angabe.

b) *DURCHGEHENDES SPRECHEN*: Die Person spricht *KONTINUIERLICH* und *OHNE UNTERBRECHUNG* von Sekunde 0 bis zur letzten Sekunde — das wird explizit und mehrfach betont, auch in jeder Zeitphase erneut ("still talking non-stop"), und endet mit "still speaking mid-sentence". Natürliche, durchgehende Lippenbewegungen.

c) *GESTIK-REFERENZ* (eigener Block, markiert mit "GESTURE REFERENCE — IMPORTANT"): Das Videomodell muss sich *EXAKT* an der Gestik der Person orientieren, so wie sie im Referenzvideo ist. Formuliere das unmissverständlich: Gesten, Hand-, Arm-, Kopfbewegungen, Körpersprache und Timing müssen dem Referenzvideo *EINS-ZU-EINS* entsprechen ("EXACTLY match", "one-to-one", "same movements at the same moments", "same rhythm and energy"). Verbiете explizit das Erfinden neuer Gesten ("Do not invent new gestures — transfer the motion from the reference video precisely onto the person in the scene"). Die Person bleibt an ihrer Position und verlässt nie den Frame.

d) *HINTERGRUND-ESKALATION*: Teile die Gesamtlänge in 3 Zeitfenster (Anfang / Mitte / Ende, z. B. bei 13 s: 0–5 / 5–9 / 9–13) und beschreibe pro Fenster, wie sich das Hintergrundereignis kontinuierlich nähert und steigert — von "klar sichtbar, aber weit weg" über "beschleunigt, wird deutlich größer, erste Partikel/Effekte erreichen die Person" bis "gefährlich nah, Effekte fliegen an der Kamera vorbei". In jedem Fenster erneut erwähnen, dass die Person weiterspricht.

e) *ENDZUSTAND*: Formuliere den gewünschten Endzustand explizit, inklusive einer klaren Verneinung dessen, was *NICHT* passieren darf (z. B. "but it does NOT engulf him"), damit das Modell die Person am Ende nicht verschluckt.

f) *STIL & VERHALTEN*: Photorealistic, cinematic, passende Licht-/Atmosphärenbeschreibung. Die Person bleibt die ganze Zeit völlig ruhig und unbeeindruckt, reagiert nicht auf das Ereignis und redet einfach lässig weiter.

REGELN:

- Der finale Prompt ist ein einziger zusammenhängender englischer Text, direkt kopierbar.

- Keine Platzhalter im finalen Prompt — alle Angaben aus Schritt 1 sind eingesetzt.
- Referenziere Startbild und Referenzvideo im Prompt exakt mit den von mir genannten Namen in Anführungszeichen.
- Gib nach dem Prompt eine kurze Stichpunkt-Erklärung (max. 4 Punkte), warum der Prompt so aufgebaut ist.

7) Prompt 3: Claude sagen, dass das Video erstellt werden soll

Erstelle jetzt das Video anhand des Prompts, dem Startbild und dem Referenzvideo. Siehe Anhang. Frage das Bild und das Video separat für Highfield ab, damit es keine Probleme gibt

8) Fertig ist das Ergebnis