



DOTCHECK

Model Card

A short summary of our in-house scoring models

VERSION	Docs v2026.7
DATE	26 July 2026
ENGINES	Vermeer · Valla · Muybridge · Helmholtz

The figures in this document come from held-out test sets. They estimate how AI-like content looks under our models — they are not proof of authenticity.

Petr Jaroš (trading as DotCheck), IČO 03330389, Jenikovice 2, 535 01 Prelouc, Czech Republic

dotcheck.ai/docs

AT A GLANCE

Images — on our held-out tests, real photos average **0.000** and AI images average **0.938** (balanced accuracy **0.9950**).

Text (English) — human writing averages **0.045** and AI writing averages **0.909** (balanced accuracy **0.980**).

Text also for Spanish, Portuguese, French, Italian, German, and Dutch — each head clears the same pass floors as English; results in the language table below.

Video frames — bag p90 real average **0.004** and AI average **0.940** (balanced accuracy **0.985**).

Audio — real average **0.004** and AI average **0.984** (balanced accuracy **0.986**).

About this Model Card

This is a short public summary of the models that power DotCheck scores. It is for customers, partners, and anyone who wants the measured results without reading the full Technical Report.

DotCheck is a Chrome extension and a website. It estimates how AI-like an image, a video (frames, and soundtrack when present), standalone audio, or text in a supported language appears. Everyday scoring uses our own models on our servers. Free, Premium, and Pro all use the same models.

What a score means

A DotCheck score is a number between 0 and 1. Higher means the content looks more AI-generated under our model for that type of media. Soft labels such as “likely human” or “likely AI” are simply a reading of that number — not a legal finding.

Treat every score as an **estimate**. It is not proof of authorship, not a courtroom deepfake opinion, and not a substitute for your own judgment.

The models

Public names (Vermeer, Valla, Muybridge, Helmholtz) sit beside stable wire ids used in the API and cache.

MEDIA	ENGINE	WHAT IT USES
Image	Vermeer (inhouse@12)	SigLIP2 vision backbone + our trained head
Text EN	Valla (inhouse-text@9)	TMR + Fakespot + English logistic head
Text ES...NL	Valla (six lang heads)	Same TMR + Fakespot bases; per-language heads
Video	Muybridge (inhouse-video@2)	SigLIP2 with a head trained for video frames
Audio	Helmholtz (inhouse-audio@2)	Dasheng-Base + logistic head; video soundtrack fuse

There is no second “better” model behind a menu. Free and Premium never send your content to Sightengine or Winston. On Pro, you can optionally ask for a commercial recheck after you already have an in-house score.

Image results

We evaluate on held-out real photos (Commons-style) and held-out AI images from Kandinsky 2.2 — a generator family that was not used to train the head. About 200 images in each group.

WHAT WE MEASURED	TARGET	RESULT	PASS
Average score on real photos	≤ 0.12	0.000	YES
Average score on AI images	≥ 0.88	0.938	YES
Gap between AI and real averages	≥ 0.55	0.970	YES
Balanced accuracy at our threshold	≥ 0.92	0.9950	YES

Text results

We score English, Spanish, Portuguese, French, Italian, German, and Dutch. Unsupported languages are not scored. Feature models are shared; each language has its own logistic head. English is the public headline claim; other languages clear the same absolute floors on their own holdouts (200 human / 200 AI each). We have not published holdouts built from live ChatGPT or Claude scrapes; generalization to those closed models is unmeasured here.

English — Valla (inhouse-text@9)

WHAT WE MEASURED	TARGET	RESULT	PASS
Average score on human text	≤ 0.12	0.045	YES
Average score on AI text	≥ 0.85	0.909	YES
Balanced accuracy at our threshold	≥ 0.90	0.980	YES

Other supported languages

Human samples are Wikipedia lead prose in each language; AI samples use open models prompted for encyclopedic prose in the same languages. Portuguese requests (pt) use the Brazilian-trained text model (inhouse-text-pt_BR_v1). All language heads share the Valla brand.

LANG	ENGINE	HUMAN MEAN	AI MEAN	BAL. ACC.
Spanish	Valla (inhouse-text-es_v2)	0.088	0.953	0.972
Portuguese	Valla (inhouse-text-pt_BR_v1)	0.103	0.945	0.928
French	Valla (inhouse-text-fr_v2)	0.082	0.976	0.975
Italian	Valla (inhouse-text-it_v1)	0.069	0.986	0.972
German	Valla (inhouse-text-de_v1)	0.064	0.968	0.958
Dutch	Valla (inhouse-text-nl_v1)	0.101	0.978	0.933

All six passed the same gates as English (TEXT_GATES_OK). Cite these measured numbers only.

Video-frame results

Short video is scored as up to three frames (VS1), then aggregated with a bag p90 — not as a full-length forensic timeline. When a soundtrack is present, product fuse with Helmholtz is separate from this bag-p90 table. Training used open video generators (CogVideoX-2b and Wan2.1). The test set used Wan2.2 clips kept out of training, plus real Commons video frames. About 100 real and 100 AI frame rows (bagged).

WHAT WE MEASURED	TARGET	RESULT	PASS
Bag p90 average on real frames	≤ 0.12	0.004	YES
Bag p90 average on AI frames	≥ 0.85	0.940	YES
Balanced accuracy at our threshold	≥ 0.90	0.985	YES

Audio results

Standalone audio (speech, music, general sound) uses Helmholtz (inhouse-audio@2): frozen Dasheng-Base + logistic head. Protocol: short windows with a max aggregate. Holdout mixes CodecFake / DFADD reals and AudioGen-style AI. Dashboard, Check, and extension-popup uploads score audio alone; video with a soundtrack may fuse window scores with the frame bag via Covenant.

WHAT WE MEASURED	TARGET	RESULT	PASS
Average score on real audio	≤ 0.12	0.004	YES
Average score on AI audio	≥ 0.85	0.984	YES
Balanced accuracy at our threshold	≥ 0.90	0.986	YES

Honest limits

- Results are for the held-out sets described above. New generators, heavy edits, or unusual compression can change scores in the wild.
- Text supports English, Spanish, Portuguese, French, Italian, German, and Dutch. We have not published holdouts from live ChatGPT or Claude scrapes; generalization to those closed models is unmeasured here.
- Video uses a handful of frames, not a full temporal deepfake timeline. Soundtrack fuse uses Helmholtz when audio is present; silent clips stay on frames.
- Audio and fused video scores are estimates for AI-likeness — not deepfake courtroom forensics.
- Optional Pro rechecks from Sightengine or Winston are separate commercial opinions. They are not part of the in-house tables above.
- A score is not enough on its own for high-stakes automated decisions about a person.

How scoring is served

Models run on DotCheck-operated servers. You do not download the weights as the normal product path. Live heads are also published under Apache-2.0 on Hugging Face. On a typical mid-range CPU host, an image score is often on the order of about 200 milliseconds; exact time depends on the machine.

Leviathan is the everyday scoring path: these models, optional generator-surface context from the Chrome extension on known AI studios, and a shared content-hash memory that can reuse a prior score for the same content later. That memory is for the content itself, not a browsing trail. It does not change the gate tables above. The public AI index uses model probability only. More detail is in the Technical Report on the same docs page.

Version

Docs v2026.7 · 26 July 2026 · Companion document: Technical Report · <https://dotcheck.ai/docs>

Upstream models (lockstep)

Commercial-clean serve backbones and train/holdout generators. SSOT: `Data.json` → `models`.

MODEL	ROLE	LICENSE
google/siglip2-base-patch16-224	image_video	Apache-2.0
Oxidane/tmr-ai-text-detector	text	MIT
fakespot-ai/roberta-base-ai-text-detection-v1	text	Apache-2.0
mispeech/dasheng-base	audio	Apache-2.0
stabilityai/sd-turbo	image fit	MIT
segmind/tiny-sd	image fit	Apache-2.0
stabilityai/sd-xl-turbo	image fit	MIT
warp-ai/wuerstchen	image fit	MIT
Qwen/Qwen2.5-7B-Instruct	text fit	Apache-2.0
mistralai/Mistral-7B-Instruct-v0.3	text fit	Apache-2.0
zai-org/CogVideoX-2b	video fit	Apache-2.0
Wan-AI/Wan2.1-T2V-1.3B-Diffusers	video fit	Apache-2.0
kandinsky-community/kandinsky-2-2-decoder	image holdout	Apache-2.0
Qwen/Qwen2.5-1.5B-Instruct	text holdout	Apache-2.0
Wan-AI/Wan2.1-T2V-1.3B-Diffusers	video holdout	Apache-2.0
Wan-AI/Wan2.2-TI2V-5B-Diffusers	video holdout	Apache-2.0

Hub cards: <https://huggingface.co/DotCheck> · <https://huggingface.co/collections/DotCheck/dotcheck-engines-6a659e57f7934cb510677785>