

Acceptable usage policy

Last Updated: September 8, 2026

This Acceptable Usage Policy ("AUP") defines the strictly prohibited boundaries, behavioral limits, and content restrictions governing your operational interaction with the artificial intelligence (AI) model aggregation platform, system, and infrastructure (collectively, the "Service") provisioned by IQBID Helios Limited ("Company," "we," "us," or "our").

This AUP is integrated directly into, and forms an inseparable part of, our European Union Terms of Service. It applies directly to all users accessing the Service from the European Union (EU) or European Economic Area (EEA) ("EU Users"). Because our system routes transactional inputs via API to third-party infrastructure entities (including OpenAI, LLC and Google LLC), you are statutory-mandated to strictly respect both our corporate standards and the downstream application rules enforced by those underlying ~~MM~~ networks.

1. Absolute prohibitions under the eu ai act and dsa

In strict compliance with the EU Artificial Intelligence Act (Regulation (EU) 2024/1689) and the Digital Services Act (Regulation (EU) 2022/2065), you are explicitly prohibited from submitting text queries, semantic instructions, prompts, or uploading any digital files (documents, images, media) (collectively, "User Inputs") designed to execute or generate the following:

A. Subliminal Manipulation and Behavioral Distortion

- **Cognitive Manipulation:** Actively prompting or forcing AI models to generate text, audio scripts, or psychological frameworks that deploy subliminal techniques beyond a person's consciousness to materially distort a person's behavior in a manner that causes or is likely to cause physical or psychological harm.
- **Exploitation of Vulnerabilities:** Generating targeted outputs designed to exploit any specific vulnerabilities of a specific group of persons due to their age, physical, or mental disability, to distort their baseline real-world actions.

B. Prohibited Biometric Systems and Social Scoring

- **Biometric Categorization:** Uploading photographic directories or facial data to classify natural persons individually based on protected characteristics (e.g., race, political opinions, trade union membership, religious beliefs, sexual orientation).
- **Untrustworthy Social Scoring:** Requesting models to evaluate, score, or categorize natural persons over a certain period based on their social behavior or known personality characteristics, leading to detrimental or unfavorable treatment in real-world

environments.

C. Illegal European Content, Violence, and Child Safety

- **Illegal Content under the DSA:** Submitting or requesting any material that constitutes illegal hate speech, xenophobic manifestos, terrorist propaganda, or operational tutorials on fabricating illicit weapons, explosives, or chemical agents.
- **Child Sexual Abuse Material (CSAM):** The Company enforces an absolute zero-tolerance threshold regarding Child Sexual Abuse Material or Child Sexual Exploitation and Abuse (CSAE). Any structural detection of material depicting minors in an exploitative or sexually explicit context will result in a permanent infrastructure ban and mandatory reporting to European law enforcement bodies (Europol/national offices).

D. Systemic Deepfakes, Deception, and Copyright Infringement

- **Deceptive Media (Deepfakes):** Generating or mutating imagery, audio, or video files that appear authentic but are synthetic ("Deepfakes") without clearly, explicitly, and permanently tagging or embedding machine-readable watermarks indicating that the content has been artificially generated or manipulated, as mandated by the EU AI Act.
- **API Harvesting and Reverse Engineering:** Utilizing systemic scraping macros, extraction web-bots, or automated loaders to extract the synthetic algorithmic training weight arrays, specific prompt configurations, or operational boundaries of our internal routing layers.

2. Moderation conduits and compliance complaints

To satisfy our statutory infrastructure duties under the Digital Services Act:

- **Automated Content Classification:** User Inputs and AI-generated Outputs are dynamically scanned by automated safety classifiers to identify high-probability systemic policy breaches before transmission across external provider nodes.
- **Notice and Action Mechanism:** If you identify any content hosted on or generated through our Service that violates local EU Member State laws or this AUP, you possess a legal right to submit a formal notification to our security department at abuse@iqbid.ai.

3. Structural account termination and financial penalties

If the Company flags clear indicators that your account interaction routines violate the mandates of this AUP, or present an immediate threat of triggering API key revocations from our primary upstream suppliers (OpenAI, Google, etc.), we will execute decisive technical actions.

CONSEQUENCES OF MATERIAL BREACH:

1. **Instant Infrastructure Revocation:** Permanent de-activation of your authorization profile, workspace partitions, and stored cache history panels without prior notification.
 2. **Cancellation of balances under the agreement:** All active subscription days for prepayment, quotas for tokens and credit units processed using the Payment Service **will be permanently canceled as an automatic compensation fee provided for in the agreement to compensate** for administrative costs for moderation.
 3. **Regulatory Disclosure:** The Company will actively cooperate with competent European data protection authorities and national markets supervisors if a formal inquiry is executed regarding illegal generative workflows.
-

Revision #4

Created 2026-09-08 14:23:36 UTC by Hostmaster

Updated 2026-09-09 11:21:11 UTC by Hostmaster