

Acceptable Use Policy

Effective Date: April 15, 2026
mAlvn, LLC | 3368 100th St, Urbandale, IA 50322

1. Overview

This Acceptable Use Policy ("AUP") governs your use of the mAlvn Developer Platform (the "Service") provided by mAlvn, LLC. This AUP supplements our Terms of Service and is incorporated therein by reference.

The Service orchestrates AI agents across multiple large language model (LLM) providers. Your use of the Service is subject not only to this AUP but also to the acceptable use policies, terms of service, and usage policies of each upstream LLM provider whose models you access through the Service. Violations of upstream provider policies may result in immediate suspension or termination of your access.

2. Upstream LLM Provider Policies

The Service currently integrates with the following LLM providers. You are responsible for reviewing and complying with each provider's usage policies:

2.1 OpenAI - GPT-4.1 series, GPT-5 series, o3, o4-mini, Codex models. Subject to the OpenAI Usage Policies.

2.2 Anthropic - Claude Haiku 4.5, Claude Sonnet 4/4.5/4.6, Claude Opus 4.5/4.6. Subject to the Anthropic Acceptable Use Policy.

2.3 Google (Gemini) - Gemini 2.0/2.5/3/3.1 Flash and Pro series. Subject to the Google Gemini API Terms and Generative AI Prohibited Use Policy.

2.4 Mistral AI - Mistral Small 3.1, Medium 3, Large, Codestral. Subject to the Mistral AI Terms of Use.

2.5 xAI (Grok) - Grok 3 Mini, Grok 4 series, Grok 4.1 series, Grok 4.20 Beta, Grok Code Fast. Subject to the xAI Terms of Service.

2.6 Future Providers and Models. The providers and models listed above reflect the Service's current integrations. We may add, remove, or replace LLM providers and models at any time without amending this AUP. When you access any model through the Service, you are automatically bound by that model's provider usage policies, whether or not the provider is explicitly named in this document. This includes providers added after the effective date of this AUP.

It is your responsibility to identify which LLM providers power the models you use through the Service, to review their respective usage policies, and to monitor those policies for changes. mAlvn is not responsible for notifying you of changes to upstream provider policies.

If an upstream provider's usage policies conflict with this AUP, the more restrictive policy applies. If you are uncertain whether a specific use case is permitted, you should refrain from that use or contact us

for guidance.

3. Universally Prohibited Uses

Regardless of which LLM provider or model you use, the following are strictly prohibited:

3.1 Illegal Activities: Violating laws, facilitating criminal activity, generating CSAM, violating export controls.

3.2 Weapons and Violence: Developing weapons of mass destruction, creating weapons instructions, promoting violence or terrorism.

3.3 Deception and Manipulation: Generating deepfakes, impersonating individuals, conducting disinformation campaigns, phishing.

3.4 Privacy and Surveillance: Unauthorized surveillance, facial recognition without consent, deanonymizing personal data.

3.5 Harm and Abuse: Promoting self-harm, harassment, hate speech, non-consensual intimate imagery, unqualified professional advice.

3.6 System Integrity: Circumventing safety measures, prompt injection, jailbreaking, reverse-engineering models, malware, denial-of-service.

3.7 High-Risk Applications: Using AI as sole decision-maker for legal/financial/safety matters, deploying in safety-critical systems without oversight, automated financial transactions without authorization.

4. AI Agent-Specific Requirements

- **Transparency:** Disclose to end users when they interact with AI agents;
- **Human Oversight:** Implement oversight for consequential agent actions;
- **Scope Limitation:** Configure agents with appropriate scope limitations;
- **Error Handling:** Implement fallback mechanisms for agent failures;
- **Data Minimization:** Agents should access only minimum necessary data;
- **Audit Trail:** Maintain records of agent actions for accountability;
- **Testing:** Adequately test agents before production deployment.

5. Content and Output Responsibilities

- Review AI-generated outputs before publishing or relying on them;
- Do not represent AI-generated content as human-created when deceptive;
- Ensure AI-generated content complies with intellectual property laws;
- You are responsible for accuracy of AI-generated content you distribute;
- Do not use the Service to generate content violating platform policies.

6. Rate Limits and Fair Use

- Comply with all rate limits and usage quotas for your subscription plan;
- Do not circumvent limits through multiple accounts or key rotation;
- Do not disproportionately consume shared resources;
- Abusive automated usage patterns are prohibited.

7. Reporting Violations

Report violations to hello@maivn.io. We investigate all reports and take appropriate action.

8. Enforcement

- Warnings for minor or first-time violations;
- Temporary suspension of access to specific models or features;
- Permanent suspension or termination of your account;
- Removal of violating content;
- Reporting to law enforcement or upstream providers as appropriate;
- Pursuing legal remedies for harmful violations.

Upstream LLM providers may also independently enforce their policies.

9. Changes to This Policy

We may update this AUP to reflect changes in upstream provider policies, laws, or our practices. We recommend periodically reviewing upstream provider policies.

10. Contact Information

For questions about this AUP or to report violations:

mAlvn, LLC
3368 100th St
Urbandale, IA 50322
United States
Email: hello@maivn.io