

Usage Policy

Effective September 15, 2025**Previous Version**
English

Our Usage Policy (also referred to as our “Acceptable Use Policy” or “AUP”) applies to anyone who can submit inputs to Anthropic’s products and/or services, including via any authorized resellers or passthrough access, all of whom we refer to as “users.” The Usage Policy is intended to help our users stay safe and promote the responsible use of our products and services.

The Usage Policy is categorized according to who can use our products and for what purposes. We will update our policy as our technology and the associated risks evolve or as we learn about unanticipated risks.

- **Universal Usage Standards:** Our Universal Usage Standards apply to all users and use cases.
- **High-Risk Use Case Requirements:** Our High-Risk Use Case Requirements apply to specific consumer-facing use cases that pose an elevated risk of harm.
- **Additional Use Case Guidelines:** Our Additional Use Case Guidelines apply to certain other use cases, including consumer-facing chatbots, products serving minors, agentic use, and Model Context Protocol servers.

Anthropic’s Safeguards Team will implement detection and monitoring to enforce our Usage Policy, so please review this policy carefully before using

our products or services. If we learn that you have violated our Usage Policy, we may throttle, suspend, or terminate your access to our products and services. We may also block or modify model outputs when inputs violate our Usage Policy.

If you believe that our model outputs are potentially inaccurate, biased or harmful, please notify us at usersafety@anthropic.com, or report it directly in our product through the “report issues” thumbs down button or similar feedback features (where available). You can read more about our Safeguards practices and recommendations in our [Safeguards Support Center](#).

This Usage Policy is calibrated to strike an optimal balance between enabling beneficial uses and mitigating potential harms. Anthropic may enter into contracts with certain governmental customers that tailor use restrictions to that customer’s public mission and legal authorities if, in Anthropic’s judgment, the contractual use restrictions and applicable safeguards are adequate to mitigate the potential harms addressed by this Usage Policy.

Universal Usage Standards

Do Not Violate Applicable Laws or Engage in Illegal Activity

This includes using our products or services to:

- Acquire or exchange illegal or controlled substances
- Engage in or facilitate human trafficking or prostitution
- Infringe, misappropriate, or violate the intellectual property rights of a third party

- Violate any other applicable laws or regulations in your jurisdiction

Do Not Compromise Critical Infrastructure

This includes using our products or services to:

- Facilitate the destruction or disruption of critical infrastructure such as power grids, water treatment facilities, medical devices, telecommunication networks, or air traffic control systems
- Obtain unauthorized access to critical systems such as voting machines, healthcare databases, and financial markets
- Interfere with the operation of military bases and related infrastructure

Do Not Compromise Computer or Network Systems

This includes using our products or services to:

- Discover or exploit vulnerabilities in systems, networks, or applications without authorization of the system owner
- Gain unauthorized access to systems, networks, applications, or devices through technical attacks or social engineering
- Create or distribute malware, ransomware, or other types of malicious code
- Develop tools for denial-of-service attacks or managing botnets
- Create tools designed to intercept communications or monitor devices without authorization of the system owner
- Develop persistent access tools designed to operate below normal system security levels, including firmware modifications or hardware implants
- Create automated tools designed to compromise multiple systems at scale for malicious purposes

- Bypass security controls such as authenticated systems, endpoint protection, or monitoring tools

Do Not Compromise Computer or Network Systems

This includes using our products or services to:

- Discover or exploit vulnerabilities in systems, networks, or applications without authorization of the system owner
- Gain unauthorized access to systems, networks, applications, or devices through technical attacks or social engineering
- Create or distribute malware, ransomware, or other types of malicious code
- Develop tools for denial-of-service attacks or managing botnets
- Create tools designed to intercept communications or monitor devices without authorization of the system owner
- Develop persistent access tools designed to operate below normal system security levels, including firmware modifications or hardware implants
- Create automated tools designed to compromise multiple systems at scale for malicious purposes
- Bypass security controls such as authenticated systems, endpoint protection, or monitoring tools

Do Not Develop or Design Weapons

This includes using our products or services to:

- Produce, modify, design, or illegally acquire weapons, explosives, dangerous materials or other systems designed to cause harm to or loss of human life

- Design or develop weaponization and delivery processes for the deployment of weapons
- Circumvent regulatory controls to acquire weapons or their precursors
- Synthesize, or otherwise develop, high-yield explosives or biological, chemical, radiological, or nuclear weapons or their precursors, including modifications to evade detection or medical countermeasures

Do Not Incite Violence or Hateful Behavior

This includes using our products or services to:

- Incite, facilitate, or promote violent extremism, terrorism, or hateful behavior
- Provide material support for organizations or individuals associated with violent extremism, terrorism, or hateful behavior
- Facilitate or promote any act of violence or intimidation targeting individuals, groups, animals, or property
- Promote discriminatory practices or behaviors against individuals or groups on the basis of one or more protected attributes such as race, ethnicity, religion, national origin, gender, sexual orientation, or any other identifying trait

Do Not Compromise Privacy or Identity Rights

This includes using our products or services to:

- Violate privacy rights as defined by applicable privacy laws, such as sharing personal information without consent or accessing private data unlawfully
- Misuse, collect, solicit, or gain access without permission to private information such as non-public contact details, health data, biometric or neural data (including facial recognition), or confidential or proprietary data

- Impersonate a human by presenting results as human-generated, or using results in a manner intended to convince a natural person that they are communicating with a natural person when they are not

Do Not Compromise Children's Safety

This includes using our products or services to:

- Create, distribute, or promote child sexual abuse material ("CSAM"), including AI-generated CSAM
- Facilitate the trafficking, sextortion, or any other form of exploitation of a minor
- Facilitate minor grooming, including generating content designed to impersonate a minor
- Facilitate child abuse of any form, including instructions for how to conceal abuse
- Promote or facilitate pedophilic relationships, including via roleplay with the model
- Fetishize or sexualize minors, including in fictional settings or via roleplay with the model

Note: We define a minor or child to be any individual under the age of 18 years old, regardless of jurisdiction. When we detect CSAM (including AI-generated CSAM), or coercion or enticement of a minor to engage in sexual activities, we will report to relevant authorities.

Do Not Create Psychologically or Emotionally Harmful Content

This includes using our products or services to:

- Facilitate, promote, or glamorize any form of suicide or self-harm, including disordered eating and unhealthy or compulsive exercise

- Engage in behaviors that promote unhealthy or unattainable body image or beauty standards, such as using the model to critique anyone's body shape or size
- Shame, humiliate, intimidate, bully, harass, or celebrate the suffering of individuals
- Coordinate the harassment or intimidation of an individual or group
- Generate content depicting animal cruelty or abuse
- Promote, trivialize, or depict graphic violence or gratuitous gore, including sexual violence
- Develop a new product or service, or support an existing product or service that employs or facilitates deceptive techniques with the intent of causing emotional harm

Do Not Create or Spread Misinformation

This includes using our products or services to:

- Create or disseminate deceptive or misleading information about, or with the intention of targeting, a group, entity or person
- Create or disseminate deceptive or misleading information about laws, regulations, procedures, practices, standards established by an institution, entity or governing body
- Create or disseminate conspiratorial narratives meant to target a specific group, individual or entity
- Impersonate real entities or create fake personas to falsely attribute content or mislead others about its origin without consent or legal right
- Provide false or misleading information related to medical, health or science issues

Do Not Undermine Democratic Processes or Engage in Targeted Campaign Activities

This includes using our products or services to:

- Engage in personalized vote or campaign targeting based on individual profiles or data
- Create artificial or deceptive political movements in which the source, scale or nature of the campaign or activities is misrepresented
- Generate automated communications to public officials or voters at scale that conceal their artificial origin, or engage in systematic vote solicitation that could undermine election integrity
- Create political content designed to deceive or mislead voters, including synthetic media of political figures
- Generate or disseminate false or misleading information in political and electoral contexts, including about candidates, parties, policies, voting procedures, or election security
- Engage in political lobbying or grassroots advocacy using false or fabricated information, or create lobbying or advocacy materials containing demonstrably false claims about facts, data, or events
- Incite, glorify or facilitate the disruption of electoral or civic processes, including interference with voting systems, vote counting, or certification processes
- Create content designed to suppress voter turnout or discourage legitimate political participation through deception or intimidation

Do Not Use for Criminal Justice, Censorship, Surveillance, or Prohibited Law Enforcement Purposes

This includes using our products or services to:

- Make determinations on criminal justice applications, including making decisions about or determining eligibility for parole or sentencing

- Target or track a person's physical location, emotional state, or communication without their consent, including using our products for facial recognition, battlefield management applications or predictive policing
- Utilize models to assign scores or ratings to individuals based on an assessment of their trustworthiness or social behavior without notification or their consent
- Build or support emotional recognition systems or techniques that are used to infer emotions of a natural person, except for medical or safety reasons
- Analyze or identify specific content to censor on behalf of a government organization
- Utilize models as part of any biometric categorization system for categorizing people based on their biometric data to infer their race, political opinions, trade union membership, religious or philosophical beliefs, sex life or sexual orientation
- Utilize models as part of any law enforcement application that violates or impairs the liberty, civil liberties, or human rights of natural persons

Do Not Engage in Fraudulent, Abusive, or Predatory Practices

This includes using our products or services to:

- Facilitate the production, acquisition, or distribution of counterfeit or illicitly acquired goods
- Promote or facilitate the generation or distribution of spam
- Generate content for fraudulent activities, schemes, scams, phishing, or malware that can result in direct financial or psychological harm
- Create falsified documents including fake IDs, licenses, currency, or other government documents

- Develop, promote, or otherwise facilitate the sale or distribution of fraudulent or deceptive products
- Generate deceptive or misleading digital content such as fake reviews, comments, or media
- Engage in or facilitate multi-level marketing, pyramid schemes, or other deceptive business models that use high-pressure sales tactics or exploit participants
- Promote or facilitate payday loans, title loans, or other high-interest, short-term lending practices that exploit vulnerable individuals
- Engage in deceptive or abusive practices that exploit individuals based on age, disability or a specific social or economic situation
- Promote or facilitate the use of abusive or harassing debt collection practices
- Develop a product or support an existing service that deploys subliminal, manipulative, or deceptive techniques to distort behavior by impairing decision-making
- Engage in actions or behaviors that circumvent the guardrails or terms of other platforms or services
- Plagiarize or submit AI-assisted work without proper permission or attribution

Do Not Abuse our Platform

This includes using our products or services to:

- Coordinate malicious activity across multiple accounts to avoid detection or circumvent product guardrails or generating identical or similar inputs that otherwise violate our Usage Policy
- Utilize automation in account creation or to engage in spammy behavior

- Circumvent a ban through the use of a different account, such as the creation of a new account, use of an existing account, or providing access to a person or entity that was previously banned
- Access or facilitate account or API access to Claude to persons, entities, or users in violation of our Supported Regions Policy
- Intentionally bypass capabilities, restrictions, or guardrails established within our products for the purposes of instructing the model to produce harmful outputs (e.g., jailbreaking or prompt injection) without prior authorization from Anthropic
- Utilization of inputs and outputs to train an AI model (e.g., “model scraping” or “model distillation”) without prior authorization from Anthropic

Do Not Generate Sexually Explicit Content

This includes using our products or services to:

- Depict or request sexual intercourse or sex acts
- Generate content related to sexual fetishes or fantasies
- Facilitate, promote, or depict incest or bestiality
- Engage in erotic chats

High-Risk Use Case Requirements

Some use cases pose an elevated risk of harm because they influence domains that are vital to public welfare and social equity. For these use cases, given potential risks to individuals and consumers, we believe that relevant human expertise should be integrated and that end-users should be aware when AI has been involved in producing outputs.

As such, for the “High-Risk Use Cases” described below, we require that you implement these additional safety measures:

- Human-in-the-loop: When using our products or services to provide advice, recommendations, or in subjective decision-making directly affecting individuals or consumers, a qualified professional in that field must review the content or decision prior to dissemination or finalization. You or your organization are responsible for the accuracy and appropriateness of that information.
- Disclosure: If model outputs are presented directly to individuals or consumers, you must disclose to them that you are using AI to help produce your advice, decisions, or recommendations. This disclosure must be provided at a minimum at the beginning of each session.

“High-Risk Use Cases” include:

- Legal: Use cases related to legal interpretation, legal guidance, or decisions with legal implications
- Healthcare: Use cases related to healthcare decisions, medical diagnosis, patient care, therapy, mental health, or other medical guidance. Wellness advice (e.g., advice on sleep, stress, nutrition, exercise, etc.) does not fall under this category
- Insurance: Use cases related to health, life, property, disability, or other types of insurance underwriting, claims processing, or coverage decisions
- Finance: Use cases related to financial decisions, including investment advice, loan approvals, and determining financial eligibility or creditworthiness

- Employment and housing: Use cases related to decisions about the employability of individuals, resume screening, hiring tools, or other employment determinations or decisions regarding eligibility for housing, including leases and home loans
- Academic testing, accreditation and admissions: Use cases related to standardized testing companies that administer school admissions (including evaluating, scoring or ranking prospective students), language proficiency, or professional certification exams; agencies that evaluate and certify educational institutions
- Media or professional journalistic content: Use cases related to using our products or services to automatically generate content and publish it for external consumption

Additional Use Case Guidelines

The below use cases – regardless of whether they are High-Risk Use Cases – must comply with the additional guidance provided.

- All consumer-facing chatbots, including any external-facing or interactive AI agent, must disclose to users that they are interacting with AI rather than a human. This disclosure must be provided at a minimum at the beginning of each chat session.
- Products serving minors, including organizations providing minors with the ability to directly interact with products that incorporate our API(s), must comply with the additional guidelines outlined in our [Help Center article](#).

- Agentic use cases must still comply with the Usage Policy. We provide examples of Usage Policy prohibitions in the context of agentic use in this [Help Center article](#).
- Model Context Protocol (MCP) servers listed in our Connector Directory must comply with our [Directory Policy](#).