

JetBrains Products Acceptable Use Policy

Version 2.1, effective as of July 17, 2025

We hope you are enjoying your JetBrains products!

To keep JetBrains cloud-based products, AI models, and services (“Product(s)”) running smoothly, safely, securely, and in compliance with applicable laws and community standards, JetBrains s.r.o. (“JetBrains”, “We”, “Us”) needs to ensure that they are used fairly and not misused by their users.

This JetBrains Products Acceptable Use Policy (“Policy”) supplements the terms of service and other policies related to each Product concerned (collectively, “Terms”) and explains what We consider unacceptable use of the Product.

Additionally, We are required by applicable law to disclose how We moderate content that We host for our users. This Policy describes the types of content not allowed for hosting within the Products, the procedures through which We moderate inappropriate content, the restrictions We may impose in the case of a breach of the Policy, the consequences of repeated misuse, and the rights of JetBrains products’ users (“You”, “Yourself”) to object to our decisions.

The Policy also covers specific restrictions related to certain types of Products, in particular AI models and Products powered by AI models hosted either by Us or by external AI service providers (“AI Service Providers”).

The procedures described in this Policy, along with the types of unacceptable use, also apply to the JetBrains website. Therefore, for the purposes of this Policy, references to “Products” include the JetBrains website as well.

1. Unacceptable Use of the Products

When You use a Product, You may not use it or attempt to use it in an unacceptable way. Unacceptable use includes:

- (a) **Illegal use.** This means using the Products in ways that violate applicable laws, regulations, or government policies.
- (b) **Undermining security.** This means compromising the integrity, performance, or security of the Products, JetBrains systems, or its infrastructure. This also includes uploading or transmitting harmful components or technologies (e.g. viruses or trojans) that interact with the Products in unauthorized or illegal ways, as well as intercepting data being sent to or from the Products or any JetBrains-related network.
- (c) **Unauthorized access.** This means circumventing any security or authentication measures, including hacking our Product or gaining unauthorized access to any computer or network. We expect You to prevent unauthorized access to the Products and keep passwords and other authentication credentials secure.
- (d) **Excessive resource use.** This means imposing an unreasonable load on our infrastructure by exceeding the parameters described in the Products’ documentation, using an unreasonable amount of JetBrains resources, or using any JetBrains-related automated systems in an unreasonable way. This includes mining cryptocurrency and other automated processes relating to cryptocurrency.
- (e) **Violating community standards.** This means behavior that goes against any applicable community standards, without respect for community participation, or that is malicious in nature. This also includes inciting, engaging in, or encouraging abuse, violence, hatred, or discrimination, as well as using language that JetBrains considers malicious.
- (f) **Impersonation or misrepresentation.** This means misrepresenting Yourself or your right to use the Products, or acting in a fraudulent or deceptive manner. This also includes misrepresenting your user rights (e.g. impersonating an admin user) or impersonating another person.
- (g) **Unsolicited communication and promotions.** This means advertising, promoting other products, or sending unsolicited communication to other users, including phishing or spoofing emails.
- (h) **Unintended use of the Products.** This means accessing and/or using the Products in any way that is contrary to what is described in the Terms and the Products’ documentation. This also includes unintended use, such as unauthorized scraping; the use of any automated or programmatic method to extract data or, where applicable,

outputs from the Products; web harvesting or web data extraction (except as permitted through the API); bulk account creation; or any similar use of the Products not authorized by JetBrains.

(i) Encouraging unacceptable use. This means using or sharing content that supports or encourages others to engage in any of the unacceptable activities listed above.

(j) Reverse engineering and source code discovery. This means reverse assembling, reverse compiling, decompiling, translating, or otherwise attempting to discover the source code of the Products and/or, where applicable, underlying components of models, algorithms, and systems of the large language models that are connected to the Products (except to the extent that such restrictions are contrary to applicable law).

(k) Sharing inappropriate content. This means uploading, publishing, displaying, featuring, storing, or otherwise making available, within the Product, any content that JetBrains may consider inappropriate, such as is described in Section 2 below.

Sending inputs and requesting outputs within the Products to achieve any of the results described in points (a)–(k) above is also strictly prohibited.

2. Inappropriate Content

2.1. Types of inappropriate Content. When You use a Product that allows You to host your content such as code, texts, articles, comments, images, photographs, graphics, software, and designs (“Content”) on servers operated by Us, You may not share or otherwise process in the Product any Content that We consider inappropriate or that falls under the categories listed below:

(a) Content infringing intellectual property rights. This includes Content that infringes others intellectual property rights (including trade secrets, copyright, trademarks, service marks, patents, and moral rights).

(b) Privacy-violating Content. This includes any Content or functionality that compromises privacy or personal data in a manner contrary to legal regulations, for instance, by accessing personal data without authorization, processing personal data without an appropriate legal basis, or otherwise using personal data in an unauthorized or illegal way.

(c) Illegal or harmful speech. This includes Content that contains defamations, discriminations, calls or incitement to violence and/or hatred, hate speech, and other similar Content.

(d) Scams and/or fraudulent Content. This includes inauthentic accounts, inauthentic user reviews, pyramid schemes, and any other Content that promotes fraudulent schemes, scams, or phishing attempts or that makes false statements about individuals or companies with the intent of harming their reputation.

(e) Spam. This includes unsolicited commercial posts unrelated to a Product, which often appear to be written by bots. The posts We qualify as spam typically contain unrelated advertisements, links to malicious websites, and irrelevant comments.

(f) Cyberviolence. This includes any forms of cyberbullying and intimidation, cyber-harassment, cyber-incitement to hatred or violence, and non-consensual sharing of material containing deepfake or similar technology using a third party’s features.

(g) Content incompatible with this Policy and/or the Terms. This includes any Content that violates, contradicts, or otherwise conflicts with the terms of the Policy and/or the Terms.

(h) Other illegal Content. This includes Content that is contrary to applicable law, regulations, or governmental policies in ways not expressly covered by this Policy.

2.2. Reporting inappropriate Content. If You believe that certain Content hosted in a Product breaches this Policy, Terms, and/or applicable law, You can report it through this form.

3. Specific Restrictions for AI Models and Products Powered by AI Models

This Section 3 applies to the following types of Products: AI models and Products powered by AI models hosted by Us or AI Service Providers.

Some restrictions included herein are required to comply with the terms and acceptable use policies of AI Service Providers. In such cases, the language adopted here is drawn from those policies, including any territory limitations and other requirements.

3.1. Territory limitations. Due to the restrictions imposed by AI Service Providers, You may use the Products powered by AI models only in locations listed here. Usage of these Products outside these territories is not allowed.

3.2. Sector-specific restrictions. Due to the restrictions imposed by AI Service Providers, the Products cannot be used for certain industry-specific activities, namely:

(a) Activities with a high risk of physical harm and critical infrastructure. This means using the Products for an activity that has a high risk of physical harm, including weapons development, military use and warfare, espionage, use for materials or activities that are subject to the International Traffic Arms Regulations (ITAR) maintained by the United States Department of State, development, management, or operation of critical infrastructure in energy and water, and in the operation of transportation technologies, or generation of content that promotes, encourages, or depicts acts of self-harm, such as suicide, non-suicidal self-injury, and eating disorders.

(b) Activities with a high risk of economic harm. This means using the Products for activities that represent a high risk of economic harm, including multi-level marketing, gambling, payday lending, automated determinations of eligibility for credit, employment, educational institutions, or public assistance services.

(c) Governmental decision-making. This means using the Products for high-risk government decision-making, including law enforcement and criminal justice, migration, and asylum.

(d) Unauthorized practice of profession. This means using the Products for engaging in the unauthorized or unlicensed practice of any profession, such as offering tailored legal, medical, accounting, or financial advice, and misleading claims of expertise or capability made particularly in sensitive areas (e.g. health, finance, government services, or law).

(e) Health-related services. This means using the Products for the purpose of telling someone that they have or do not have a given health condition, or offering instructions on how to cure or treat a given health condition.

(f) Politics. This means using the Products for political campaigning or lobbying by generating high volumes of campaign materials, building conversational or interactive systems such as chatbots that provide information about campaigns or engage in political advocacy or lobbying, or building products for political campaigning or lobbying purposes.

(g) Adult content, dating, and sexual services. This means using the Products for adult industries, the creation of adult content (including content intended to arouse sexual excitement, such as the description of sexual activity), content that promotes sexual services (excluding sex education and wellness), erotic chat, dating apps, or pornography.

3.3. Unacceptable use of AI models and Products powered by AI models. In addition to what was described in Section 1 above, the following uses of AI models and the Products powered by AI models are also unacceptable:

(a) Violating privacy. This means using the Products for the purpose of any activity that violates any individual's privacy, such as tracking or monitoring an individual without their consent, facial recognition of private individuals, classifying individuals based on protected characteristics, using biometrics for identification or assessment, the unlawful collection or disclosure of personally identifiable information or educational, financial, or other protected records, generating personally identifying information for distribution or other harms, as well as generating content that may have unfair or adverse impacts on people, particularly impacts related to sensitive or protected characteristics.

(b) Generating violent content. This means generating extremism or terrorist content as well as hateful, harassing, or violent content, such as content that expresses, incites, or promotes hate based on identity, content that intends to harass, abuse, threaten, or bully an individual, or content that encourages to commit any type of crimes, promotes or glorifies violence, or celebrates the suffering or humiliation of others.

(c) Creating harmful code or engaging in deceptive activity. This means generating code that is designed to disrupt, damage, or gain unauthorized access to a computer system (for example, creating malware), or engaging in activities of a fraudulent, defamatory, or deceptive character, including scams, coordinated inauthentic behavior, false online engagement, plagiarism, disinformation, or spam.

- (d) Developing AI models with the AI models of the AI Service Providers. This means using the Products to develop AI models that compete with the AI models of the AI Service Providers.
- (e) Misrepresenting AI output as human work. This means representing that the output generated by or within the Product was a human-generated or original work when it was not.
- (f) Exploiting or harming children. This means sending to the Products any personal information of children under the age of 13 or the applicable age of digital consent, uploading or generating any child sexual abuse material or any content that exploits or harms children, or harming minors or interacting inappropriately with people under 13 years of age in any other way.
- (g) Promoting or facilitating illegal substances or services. This means using the Products to promote or facilitate the sale of, or provide instructions for synthesizing or accessing, illegal substances, goods, or services.
- (h) Automated decision-making in sensitive domains. This means making automated decisions in domains that affect material or individual rights or well-being (e.g. finance, law, healthcare, housing, insurance, and social welfare).

4. How We Moderate Inappropriate Content and Handle Other Breaches

4.1. Inappropriate Content moderation. We do not use any algorithmic decision-making, and all moderation activities are carried out through human review. When We discover, either based on a user report or as a result of our own moderation activities, that Content placed in a Product is illegal or incompatible with this Policy or the Terms, We decide on the restrictions to be imposed on the Content and/or the user in breach. These may include removal of the inappropriate Content; temporary or permanent, partial or complete suspension of its accessibility to other users; blocking access to such Content from certain locations; or suspension or termination of the provision of the respective Product, in whole or in part, to You or your user in breach (such suspension may include temporary suspension or permanent cancellation of the subscription or other type of access to the respective Product).

4.2. Repeated sharing of inappropriate Content. If You repeatedly breach this Policy by sharing inappropriate Content, We may suspend the provision of the respective Product (for example, your ability to share the Content within the Product or your subscription or other type of access to the respective Product) for a period of 30 days. After this period, the provision of the Product will be resumed. However, if You continue to share Content that violates the Policy, the provision of the services may be suspended permanently.

4.3. Misuse of the Content reporting mechanism. If You frequently report Content shared by another user of the Product that, upon our evaluation, is not illegal and does not breach this Policy and/or the Terms, You may receive a warning. Continued misuse of the reporting mechanism after the warning may lead to the suspension of Your right to the processing of Content reports submitted by You for 30 days. Reports submitted during this period will not be processed, even after it ends.

4.4. Misuse of the complaint-handling mechanism. If You frequently submit complaints according to Section 5(a) that, upon our evaluation, are manifestly unfounded, You may receive a warning. Continued misuse may lead to the suspension of the processing of your complaints for 30 days. Claims submitted during this period will not be processed, even after it ends.

4.5. Other breaches. With regard to other breaches – such as unacceptable use of the Products, violations of the specific restrictions for AI models and Products powered by AI models, and/or breaches of the Terms – We may also suspend or terminate the provision of the respective Product, in whole or in part, to the user in breach (for example, We may temporarily suspend or permanently cancel the subscription or other type of access to the respective Product).

4.6. No refund in case of a breach. In the event of a breach of this Policy and/or the applicable Terms, we reserve the right to take the measures described above in relation to the relevant Product, without any obligation to refund prepaid amounts.

4.7. Communication of decisions to users. Any decisions made will be duly communicated, where applicable, to all affected users, including those who submitted the initial report and/or claim.

5. How You Can Object to Our Decisions

If You disagree with our decision about any action with respect to the Content that You shared within the Products or that You reported to Us, or if You disagree with our decision to suspend or terminate the provision of the respective Product, You have the following options:

(a) Internal complaint-handling mechanism. You have the right to lodge a complaint against our decision, and We have the obligation to revisit this decision. If You want to use this right, send your complaint to dsa-complaint@jetbrains.com within six months of the date of the decision. In your complaint, include the following details:

- (i) the decision number (provided in the notification about the decision);
- (ii) a clear explanation of why You believe our decision was incorrect;
- (iii) optionally, any supporting evidence, such as screenshots or documents, that We should consider as part of the review (for example, proof of your ownership of intellectual property rights, etc.).

Upon receiving your complaint, We will send You a confirmation of receipt. We may also request additional information if necessary for the review of your complaint. After completing our review, We will inform You of the outcome and provide a justification for the decision.

(b) Out-of-court dispute settlement. You have the right to select any out-of-court dispute settlement body that has been certified under the applicable law of the European Union. Please note that these out-of-court dispute settlement bodies resolve disputes between users and online platforms only. That means You can use this procedure to resolve disputes related to any decisions We make concerning our online-platform Products, which provide public access to the Content, including the decisions about your complaints submitted through the internal complaint-handling system.

The options described above do not deprive You of your right to seek a remedy through a standard legal procedure.

6. Miscellaneous

6.1. Replacement of previous policies. From its effective date, this Policy replaces the JetBrains AI Acceptable Use Policy and Kineto Acceptable Use Policy.

6.2. Updates to the Policy. The Policy may be amended by us from time to time by publishing the new version of the Policy on the JetBrains website. We recommend that You regularly check the Policy for changes and/or additions.

In case of significant changes in the Policy, We will additionally inform You either by displaying them in your JetBrains Account or by sending the updated version of the Policy to the email address that is linked to your JetBrains Account.

Any changes to the Policy will take effect on the date specified in the updated Policy. By continuing to use the Products after the effective date, You agree to be bound by the modified Policy.