

Instagram Account Suspensions

The complete field guide to every enforcement label.

Ordered by how often they hit — with what the alert looks like at login.

DOCUMENT IG-ENF-2026-014	CLASSIFICATION Read-Only	DISTRIBUTION Advisory / restricted
REVISION 3.2 — 2026	SYSTEM OF RECORD Community Standards	COMPILED BY CherryRed PR

§ 0 ABOUT THIS REFERENCE

This guide maps every reason an Instagram (Meta) account can be restricted, suspended, or permanently disabled. Categories match Meta's **Community Standards** — the unified rulebook for Instagram, Facebook, Messenger and Threads — plus the account-security and algorithmic triggers enforced automatically.

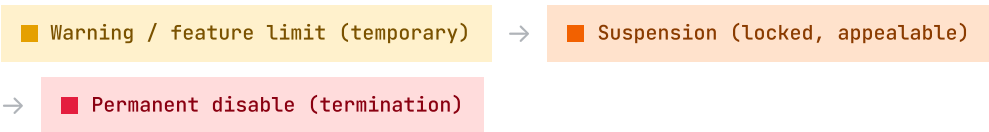
SYSTEM NOTE · ORDERING

Sections run from **most common** (Integrity, Authenticity & Automation) down to the rarer, high-severity safety categories.

SYSTEM NOTE · WORDING

Instagram rarely names the exact sub-policy in the login alert. The clearest sub-reason lives in **Settings > Account > Account Status**.

SEVERITY SCALE



CERRYRED ADVISORY — ANNOTATION ADDED BY EXTERNAL ADVISOR

Read the reinstatement section before advising any client to appeal. The appeal is effectively one-shot — spending it badly is worse than not spending it yet. If an account is still **unappealed**, that is the strongest position to recover from.

QUICK-REFERENCE CHEAT SHEET

Find your alert → see if it's appealable → send it to CherryRed **before you act.**

ENFORCEMENT LABEL	WHAT YOU'LL SEE	APPEALABLE?
Spam / Inauthentic Activity	Grey "Try Again Later" + countdown; buttons greyed out	■ Usually self-lifts
Inauthentic Behavior (IB/CIB)	"We suspended your account" / disabled for Terms	■ Yes (single); rare at scale
Account Integrity / Impersonation	"Disabled"; may ask for gov ID or video selfie	■ Yes; not ban-evasion
Cybersecurity / Suspicious Login	"Confirm it's you" security checkpoint (code / reset)	■ N/A — verify, not appeal
Misinformation	"False information" label + reduced reach	■ Yes; rarely a ban
Adult Nudity & Sexual Activity	"Removed — nudity / sexual activity"; repeat → suspend	■ Yes; common false-positive
Sexual Solicitation	Removal / suspension citing solicitation	■ Yes; enforced firmly
Violent & Graphic Content	"Sensitive Content" cover you tap through	■ Yes
Copyright	Muted audio / "copyright strike"; strikes → disabled	■ Yes; DMCA counter
Trademark	"Removed — report from IP owner"	■ Yes, with proof
Bullying & Harassment	"Removed — bullying / harassment"	■ Yes
Hateful Conduct	"Removed — hate speech"; repeat → disable	■ Yes
Adult Sexual Exploitation (NCII)	"Disabled"; victims pre-block via StopNCII	■ Yes
Restricted Goods & Services	"Removed — restricted goods"; repeat → disable	■ Yes; not at scale
Suicide / Self-Injury / Eating Disorders	Support-resource screen or removal	■ Yes; content-level
Violence & Incitement	"Removed — violence"; credible threat → disable	■ Yes; not if severe
Dangerous Orgs & Individuals	"Disabled" — designated-entity links	■ No; egregious
Human Exploitation / Trafficking	"Your account has been disabled"	■ No; egregious
Privacy Violations (Doxxing)	"Removed — shares private information"	■ Yes
Child Safety (CSE)	"Disabled" citing child sexual exploitation	■ No + — fight a false flag
Legal / Local-Law Removal	"Not available in your country" / geo-block	■ Varies by jurisdiction
Memorialization	Profile shows "Remembering"; frozen	■ N/A — not a penalty

■ appealable / recoverable ■ conditional ■ egregious — appeal usually bypassed

CHERRYRED ADVISORY — BEFORE YOU APPEAL

The standard review is effectively **one-shot**. Send the alert plus details to Justin@cherryredpr.com or the form at cherryredpr.com, ideally inside the ~30-day appeal window (the on-screen 180 days is a deletion clock, not a suspension length).

1 Integrity, Authenticity & Automation

Where most creator, performer and business accounts actually get caught. Mostly automated — fast, often no human, frequently "no reason given."

Spam / Inauthentic Activity		AI
Integrity · "We restricted some of your activity"		
SEVERITY	■ Action block / feature limit → ■ Suspension if it persists	
ENFORCED BY	Automated detection, often no human review — rate / behaviour limits enforced instantly.	
YOU'LL SEE	A grey pop-up: "Try Again Later — We limit how often you can do certain things to protect our community," often with a countdown ("Try again in 23:47"). Like / follow / comment buttons go grey.	
TRIGGERS	High-volume or automated actions: mass follow/unfollow, repetitive comments/DMs, misleading links, and third-party "growth," auto-liker, or "who unfollowed me" apps.	
EXAMPLES	Follow/unfollow bots, engagement pods, mass-DM tools, comment spam, link spam.	
APPEAL	Short blocks usually lift on their own. Disconnect third-party apps first. Persisting after warnings escalates to full suspension.	

Inauthentic Behavior (IB / CIB)		AI
Integrity · Fake engagement / fake accounts		
SEVERITY	■ Suspension → ■ Permanent at scale	
ENFORCED BY	Behavioural pattern detection, largely automated — often no human review.	
YOU'LL SEE	"We suspended your account" with a countdown, or "Your account has been disabled for not following our Terms." Usually no specific behaviour named.	
TRIGGERS	Fake accounts, artificially boosting reach, coordinated inauthentic behaviour, or misusing multiple accounts to mislead. A 2025 update spares authentic communities co-opted by CIB operations.	
EXAMPLES	Bought followers/likes, engagement pods, mass-created accounts, coordinated networks.	
APPEAL	Appealable via Account Status, but at-scale networks are rarely reinstated. Single-account false flags are the recoverable ones.	

Account Integrity & Authentic Identity	
Integrity · Impersonation / ban evasion	
SEVERITY	Suspension → Permanent
ENFORCED BY	AI flags, humans review escalations.
YOU'LL SEE	"Your account has been disabled." For impersonation you may be asked to confirm your identity with a photo of a government ID or a video selfie.
TRIGGERS	Misrepresenting who you are; impersonating a person/brand; running an account via unauthorized automation/scripting; ban evasion (new account after a prior ban); close linkage to a network of violating accounts (the "carpet ban" across shared device/IP/email).
EXAMPLES	Impersonation of a celebrity/brand, ban-evasion accounts, bot-created accounts, buying/selling account access.
APPEAL	Impersonation appeals may require ID or proof of authorization. Ban evasion is treated as egregious — keep personal and brand accounts on separate device / SIM / email / Wi-Fi.

Cybersecurity / Suspicious Login	
Security · "Suspicious login attempt"	
SEVERITY	Temporary Lock (checkpoint)
ENFORCED BY	Security systems, automated — often no human review. Not a policy strike.
YOU'LL SEE	"We detected unusual activity" / "Help us confirm it's you" — a security checkpoint: enter a code, reset your password, or confirm recent logins.
TRIGGERS	Phishing, unauthorized-access attempts, malware/malicious links, a compromised account, or simply logging in from new devices/locations/VPNs or in unusual bursts.
EXAMPLES	Hacked account sending spam, credential-harvesting, bot-like login bursts, new-country login.
APPEAL	Resolve via the checkpoint (verify identity, reset password). If hacked, use instagram.com/hacked — NOT the policy-appeal flow.

Misinformation Integrity · False information		BOTH
SEVERITY	■ Label / reduced reach → ■ Escalation for repeat harm	
ENFORCED BY	AI flags, humans review escalations.	
YOU'LL SEE	A label over the post ("False information" / context screen) and quieter reach, rather than a lock. Rarely a standalone account ban.	
TRIGGERS	Content likely to cause specific harm (some health, electoral / voter-suppression, manipulated media). Meta ended US third-party fact-checking in 2025 (moving toward Community Notes), so demotion is now more common than removal.	
EXAMPLES	Voter-suppression falsehoods, harmful health hoaxes, some manipulated media.	
APPEAL	Rarely a ban unless harmful and repeated; usually labelling or demotion.	

2 Nudity & Sexual Content

High-relevance for adult-content accounts. Heavily AI-driven image classifiers — a major source of false positives.

Adult Nudity & Sexual Activity Community Standards · Nudity		AI
SEVERITY	■ Content removal → ■ Suspension for repeat	
ENFORCED BY	Image classifiers do most detection; human review on appeal.	
YOU'LL SEE	"We removed your post — it goes against our Community Standards on nudity or sexual activity." Repeat removals escalate to "We suspended your account." Content may also just be age-gated / hidden.	
TRIGGERS	Real or AI-generated nudity, intercourse, or explicit sexual content. Some artistic / health / protest nudity is allowed. Female nipples restricted with documented exceptions (mastectomy, breastfeeding, protest).	
EXAMPLES	Explicit photos, pornographic content, digitally generated nudity, uncovered genitalia.	
APPEAL	Appealable and a common false-positive source (art, breastfeeding, medical). Cite the specific exception in the appeal.	

Sexual Solicitation Community Standards • Sexual solicitation BOTH	
SEVERITY	Suspension
ENFORCED BY	AI flags, humans review escalations.
YOU'LL SEE	"We suspended your account" or post removal citing solicitation. Often triggered by explicit offers or "menu"-style pricing in captions/DMs, or off-platform sex-work links.
TRIGGERS	Offering or asking for paid sexual services, or coordinating sexual encounters via explicit sexual language. This is the line adult accounts cross most often without realising — implicit is tolerated, explicit solicitation is not.
EXAMPLES	Escort / "selling content" offers with explicit terms, pricing for sexual acts, solicitation in DMs.
APPEAL	Appealable, but commercial sexual solicitation is enforced firmly. Keeping captions suggestive-not-explicit and links compliant is the prevention lever.

Violent & Graphic Content Community Standards • Graphic content AI	
SEVERITY	Warning screen / age-gate → Escalation for glorification
ENFORCED BY	AI first, human on appeal.
YOU'LL SEE	A "Sensitive Content" cover you tap through, and hidden-from-under-18. Escalates to removal / suspension only for glorifying content.
TRIGGERS	Gore or graphic imagery that glorifies violence or celebrates suffering. Awareness / newsworthy graphic content may stay behind a warning screen.
EXAMPLES	Glorifying a beating, celebrating an injury, sadistic gore. (Awareness footage may be allowed with a screen.)
APPEAL	Appealable; the glorification-vs-awareness distinction is the deciding factor.

3 Intellectual Property

Complaint-driven and among the most common causes of suspension. A repeat-infringer count builds toward termination.

Copyright

■ HUMAN

Copyright • the label you literally see

SEVERITY

■ Strike / removal

→

■ Permanent under repeat-infringer policy

ENFORCED BY

Report- or complaint-driven (DMCA notices), plus AI (Rights Manager / audio matching).

YOU'LL SEE

"Your post was removed because it may contain [song/owner]'s content," or muted audio, or a copyright strike notice in Account Status. Multiple strikes → "Your account has been disabled."

TRIGGERS

Posting music, video, images, or text you don't own or have a license for. A single valid DMCA can trigger action if content stays up after warning; accumulated valid takedowns lead to termination.

EXAMPLES

Unlicensed music in a Reel, reposting others' videos, using copyrighted images.

APPEAL

File a DMCA counter-notification if the claim is wrong. Removing flagged content promptly prevents escalation to account level.

Trademark

■ HUMAN

Intellectual Property • Trademark

SEVERITY

■ Removal

→

■ Suspension

ENFORCED BY

Rightsholder complaint-driven.

YOU'LL SEE

"Your content was removed in response to a report from the intellectual property owner." Often targets brand names, logos, or a storefront's trade dress.

TRIGGERS

Using a protected brand name, logo, or trade dress in a way that misleads people about source or affiliation. One well-supported complaint can trigger action.

EXAMPLES

Fake brand storefront, counterfeit sales, misleading use of a company's logo.

APPEAL

Appealable with proof of authorization, license, or fair / nominative use.

4 Harassment, Hate & Human Dignity

Mostly human- and report-driven, so context matters. Adult accounts are often on the receiving end — targeted reporting can trigger these.

Bullying & Harassment	
Community Standards · Bullying	
■ HUMAN	
SEVERITY	■ Feature limit → ■ Suspension for targeted / repeat abuse
ENFORCED BY	Mostly report-driven, context-specific.
YOU'LL SEE	"We removed your comment/post — bullying or harassment." Repeat / targeted behaviour → suspension. Victims see a report-outcome notice in Support Requests.
TRIGGERS	Repeated unwanted contact, targeted insults, sexual harassment, threats, or degrading content aimed at an individual. Private individuals get more protection than public figures.
EXAMPLES	Coordinated pile-ons, sexualized comments, mass-DM harassment, content to demean someone.
APPEAL	Appealable. Because it's context-specific, detailing the targeting in the appeal helps. Coordinated false-reporting of a performer can also be reported back.

Hateful Conduct	
Community Standards · Hate speech	
■ BOTH	
SEVERITY	■ Removal → ■ Suspension → ■ Permanent for repeat
ENFORCED BY	AI flags, humans review escalations.
YOU'LL SEE	"We removed your content — hate speech." Repeat → suspension / disable.
TRIGGERS	Direct attacks or dehumanizing speech against people based on protected characteristics (race, ethnicity, national origin, religion, caste, sex, gender identity, sexual orientation, disability, serious disease). Definition narrowed Jan 2025 to focus on direct attacks.
EXAMPLES	Slurs, dehumanizing comparisons, calls for exclusion based on a protected trait.
APPEAL	Appealable. Enforcement now centers on direct attacks rather than broad topical discussion.

Adult Sexual Exploitation (NCII) Community Standards · Sexual exploitation		BOTH
SEVERITY	Suspension → Permanent	
ENFORCED BY	AI flags, humans review escalations.	
YOU'LL SEE	"Your account has been disabled." Victims of leaked images can pre-block via StopNCII.	
TRIGGERS	Non-consensual sexual content, threats to share intimate images ("sextortion"), or sharing of non-consensual intimate imagery (NCII). Especially relevant when a performer's paid content is leaked / reposted.	
EXAMPLES	Revenge porn, threatening to leak nudes, sharing hidden-camera or stolen paid content.	
APPEAL	Appealable; performers can also file NCII / StopNCII reports to remove stolen content from the platform.	

5 Safety & Serious Harm

Rarest for typical accounts but the most severe — several skip straight to permanent, non-appealable disable.

Restricted Goods & Services Community Standards · Restricted goods		BOTH
SEVERITY	Suspension → Permanent for repeat / at-scale	
ENFORCED BY	AI flags, humans review escalations.	
YOU'LL SEE	"We removed your post — it goes against our standards on restricted goods." Repeat → disable.	
TRIGGERS	Attempting to buy, sell, trade, or gift non-medical drugs, pharmaceuticals, firearms, or other regulated goods; coordinating theft or fraud.	
EXAMPLES	Selling weed / pills, gun sales, hacking-for-hire, coordinating theft.	
APPEAL	Appealable; commercial-scale or repeat offenders are less likely to be reinstated.	

Suicide, Self-Injury & Eating Disorders	
Community Standards · Sensitive content	
■ AI	
SEVERITY	■ Warning / removal → ■ Escalation if promoting / encouraging
ENFORCED BY	Automated detection, plus human review for context / crisis referral.
YOU'LL SEE	A support-resource screen ("Can we help?") over content, or removal. Promotion escalates to suspension.
TRIGGERS	Content promoting, encouraging, or instructing suicide, self-harm, or disordered eating. Awareness / recovery content is allowed but may be age-gated or covered by a warning.
EXAMPLES	"Thinspo" / pro-ED content, self-harm method instructions, glorification of suicide.
APPEAL	Content usually removed rather than the account disabled on first instance; repeated promotion escalates. A sensitive area — if a client is personally affected, direct them to real support, not just appeals.

Violence & Incitement	
Community Standards · Violence	
■ BOTH	
SEVERITY	■ Suspension → ■ Permanent for credible threats
ENFORCED BY	AI flags, humans review escalations. Meta coordinates with law enforcement for genuine risk.
YOU'LL SEE	"We removed your content — violence or incitement," or immediate disable for credible threats.
TRIGGERS	Threats of violence, statements of intent, calls to harm, or content facilitating serious offline harm.
EXAMPLES	Credible death threats, calls to bring weapons to a location, coordinating violence.
APPEAL	Appealable unless there is an extreme safety concern, in which case the standard appeal is bypassed.

Dangerous Organizations & Individuals

■ BOTH

Community Standards · Terrorism / Organized hate

SEVERITY ■ Permanent disable

ENFORCED BY AI flags, humans review escalations.

YOU'LL SEE "Your account has been disabled" — typically no appeal path for designated-entity links.

TRIGGERS Praise, support, or representation of designated terrorist groups, organized hate groups, or criminal orgs, or accounts run on their behalf.

EXAMPLES Terrorist propaganda, running an account for a banned org, glorifying a mass shooter.

APPEAL Very limited; treated as egregious.

Human Exploitation & Trafficking

■ BOTH

Community Standards · Human exploitation

SEVERITY ■ Permanent disable

ENFORCED BY AI flags, humans review escalations.

YOU'LL SEE "Your account has been disabled." Egregious — minimal appeal.

TRIGGERS Facilitating human trafficking, smuggling, forced labor, or other exploitation of people.

EXAMPLES Recruiting for trafficking, advertising smuggling, forced-labor coordination.

APPEAL Treated as egregious; limited appeal.

Privacy Violations (Doxxing)

■ HUMAN

Community Standards · Privacy

SEVERITY ■ Removal → ■ Suspension

ENFORCED BY Report-driven.

YOU'LL SEE "We removed your content — it shares private information."

TRIGGERS Sharing others' private / personally identifiable information without consent, or content that violates someone's privacy.

EXAMPLES Posting someone's home address, phone number, ID documents, or private medical info.

APPEAL Appealable; the affected person can also file a privacy report.

Child Safety (CSE)	
Community Standards · Child safety / "child sexual exploitation"	
SEVERITY	■ Permanent disable – zero tolerance
ENFORCED BY	Automated detection + specialized human teams + PhotoDNA hash-matching. Reported to NCMEC / law enforcement.
YOU'LL SEE	"Your account has been disabled" citing child sexual exploitation. Often no standard appeal. Note: the 2025–26 AI waves produced many false CSE flags on adult accounts — devastating and hard to reverse.
TRIGGERS	Any sexualized content involving minors; grooming; solicitation; even certain non-sexual minor nudity. Also: off-platform child-sexual-offence convictions. Under-13 accounts are removed.
EXAMPLES	Sexualized imagery of anyone under 18; sexualized text about minors; sharing / soliciting CSAM.
APPEAL	Extremely limited — egregious accounts skip the standard appeal. For a false CSE flag on an adult account, act immediately with documented proof of content history and age of everyone depicted.

6 Legal & Account-State

Account states that look like enforcement but aren't policy strikes — handled through their own flows, not the violation-appeal path.

Legal / Local-Law Removals	
Legal · Local law	
SEVERITY	■ Varies – content restriction to account action
ENFORCED BY	Legal / policy teams.
YOU'LL SEE	"This content isn't available in your country," a geo-restriction, or a legal-request notice.
TRIGGERS	Government / court-ordered removals under local law, valid law-enforcement requests, or accounts prohibited from Meta's services under applicable law (e.g. sanctions).
EXAMPLES	Court-ordered geo-restriction, sanctioned individuals, legally mandated takedowns.
APPEAL	Depends on jurisdiction; may be a geo-restriction rather than a global ban.

Memorialization Account status • Memorialized		■ HUMAN
SEVERITY	■ Account locked / memorialized – not a punishment	
ENFORCED BY	Handled by dedicated teams on verified request.	
YOU'LL SEE	The profile shows "Remembering" and is frozen from new logins.	
TRIGGERS	Not a violation. When Meta is notified a user has died, the account can be memorialized or removed at a verified family member's request.	
EXAMPLES	Family submits proof of death; account is frozen.	
APPEAL	Handled through a dedicated memorialization request, not the violation-appeal flow.	

Why unappealed cases are easier to save

The mechanic that explains it, grounded in Meta's own policy wording.

The appeal is a one-shot, terminal review

Meta's Help Center states plainly: *"if you do not appeal within the specified time, you will no longer be able to request a review of your account."* Once the appeal is used and denied, the standard review door closes — only heavy escalation paths (regulators, press, the Oversight Board) can move it.

An unappealed account still holds its **review token unspent**, so every softer channel remains open.

Why a hasty appeal actively hurts

- Most first appeals are auto-denied in minutes by software, not a human — a weak appeal burns the token on an instant "no."
- Rapid or duplicate resubmissions are flagged as spam and correlate with slower outcomes.
- An unspent appeal keeps the good channels open: one clean evidenced appeal, Meta Verified live chat, Business/Ads support, and regulator routes.

180 DAYS

The suspension-screen countdown is a **data-deletion clock**, not a suspension length. At day 180 the account, content and username are deleted and reinstatement becomes impossible.

~30 DAYS

The clock that actually governs the appeal is much shorter. Make the strong move **well before** this window closes.

Recommended handling sequence — per client account

1	Preserve. Screenshot the exact alert (with countdown + named policy), save Meta's email, note the case number.
2	Export data. If any login still works: Accounts Center > Your information and permissions > Download your information. Often unavailable once fully disabled.
3	Diagnose the state. Suspended (countdown, limited login) → Disabled (web forms only) → Permanently disabled (escalation only). The wording tells you which options remain.
4	Route correctly BEFORE appealing. Hacked → instagram.com/hacked . Security checkpoint → verify. Copyright → DMCA counter-notice. Impersonation → ID form. Only genuine policy calls use the suspension appeal.
5	Spend the appeal once, well. One calm, specific, evidenced appeal (300–500 words; no confessions, no duplicates). This is the token — don't waste it on a vague first draft.
6	Escalate in parallel, not by re-appealing. If denied / silent past ~2 weeks: Meta Verified live chat, Business/Ads support, and (EU/UK) GDPR Art. 22 / Oversight Board — all inside the 180-day clock.

CHERRYRED ADVISORY — CONTACT BEFORE YOU APPEAL

An unappealed account still has its one review token unspent — the strongest position to recover from.

If that's you, **do not file the appeal yourself** — send it to us first so we spend that single shot correctly. Reach Justin@cherryredpr.com or the form at cherryredpr.com, and include:

- Screenshot of the exact alert (with the countdown and the named policy)
- Which state you're in — suspended (countdown showing), disabled, or permanently disabled
- The enforcement label / category it appears to fall under (from this guide)
- Whether you've appealed already — and if so, what you said and the outcome
- The suspension date and how many days the countdown shows remaining
- Any Meta email / case number, plus your username and whether the account has ever run ads

SCAM WARNING

No one can pause or reset the 180-day countdown, and no paid service can restore an account Meta has decided to keep down. Anyone promising a guaranteed “unban” for a fee — especially if they ask for the login — is a scam.

CAN YOU TELL IF YOU'RE BEING REPORTED?

Short answer: you can **never** see who reported you, and you get no alert at the moment a report is filed. Instagram keeps the reporter anonymous by design, and a single report just enters a review queue. But you can pick up indirect signals.

× WHAT YOU CAN'T SEE

- **The reporter's identity** — ever. Any app claiming to “reveal who reported you” is a scam.
- A **notification** that a report was filed — you're only told about an outcome.
- A **count** of reports against you, or your internal “trust score.”

✓ WHAT YOU CAN CHECK

- **Settings › Help › Support Requests.** Actioned reports show their outcome here — you can “Request a Review” from the same screen.
- **Account Status.** Removed content, reach limits, feature restrictions. A new entry is the clearest sign a flag landed.
- A sudden **reach drop** with no other cause — a prompt to check the dashboards above.
- **Hostile comments or DMs** right before a strike — reports often follow a visible conflict.

CHERRYRED ADVISORY — FOR PERFORMERS

One report is near-harmless — it's **coordinated mass-reporting** that does damage, because volume pushes content into human review and dents your trust score. If you suspect a brigading campaign, document the timeline (screenshots, timestamps) and send it to CherryRed.

TALL TALES VS FACTS

What holds up against Meta's actual behaviour and what doesn't. Where the truth is "it depends," we say so.

"Change your email & phone after reinstatement and you won't get suspended again."

TALL TALE

REALITY Swapping credentials on a recovered account does not lower your suspension risk. Enforcement keys off behaviour, content, reports and device/IP fingerprint — not which email is attached. A reinstated account is on a probationary period where new flags are treated more harshly, so what protects you is reducing activity for 1–2 weeks, cutting third-party apps, and not repeating the trigger.

"You need a new email & phone to make a fresh account after a permanent ban."

PARTLY TRUE

REALITY True that old credentials won't re-open a disabled account and a new account needs new ones — but new email/phone alone is not enough. Meta links accounts by device ID, IP, and Wi-Fi (the "carpet ban"). A new account on the same phone/network can be tied to the banned one. Note: making a new account to evade a ban is itself a policy violation (ban evasion).

"Just wait out the 180 days and your account comes back."

TALL TALE

REALITY The opposite. The 180 days is a data-deletion countdown, not a sentence you serve. Nothing is restored at the end — day 180 is when recovery becomes impossible because the data is wiped. The window that matters is the ~30-day appeal window inside it.

"Appealing many times — or from several forms at once — improves your odds."

TALL TALE

REALITY Rapid or duplicate appeals get flagged as spam and correlate with slower outcomes; the standard review is effectively one-shot. One clean, specific, evidenced appeal beats ten copy-pasted ones. This is exactly why an unappealed account is the best starting position.

"The "I was hacked" ([instagram.com/hacked](https://www.instagram.com/hacked)) flow is a backdoor to reverse a policy ban."

PARTLY TRUE

REALITY It can route you to a human identity-verification path, and some people report success — but using the hacked flow for a genuine policy strike is the wrong form and can be treated as bad-faith, hurting your case. Use it only if the account was actually compromised.

"Meta Verified guarantees you get a human and your account back."

TALL TALE

REALITY Verified can improve access to live support and may add trust, but it does not change review criteria or override automated enforcement. During the 2025 mass-ban wave, paying subscribers were banned alongside free accounts and often got little help. Useful, not magic.

"A VPN / new IP lets a banned user quietly slip back on."

PARTLY TRUE

REALITY An IP change alone rarely defeats detection because Meta also fingerprints the device and behaviour. A truly clean setup (different device, network, credentials) reduces linkage, but none of this makes ban evasion compliant — it remains a violation if you're evading an enforcement action.

"You can find out who reported you."

TALL TALE

REALITY Never. Reporter identity is permanently anonymous. Any tool promising to reveal it is a scam. You can only see outcomes in Support Requests / Account Status.

"Posting to Stories instead of the feed avoids enforcement."

PARTLY TRUE

REALITY Anecdotally Stories sometimes draw lighter automated scrutiny, but the same Community Standards apply to Stories, Reels and feed alike. Treat any "safe surface" claim as a soft tactic, not a rule — violating content is actionable wherever it's posted.

Sources: Meta Transparency Center (Community Standards & Account Integrity), Meta Help Center (suspended / disabled accounts), and 2025–26 recovery / enforcement reporting. Category names reflect Meta's unified Community Standards; exact on-screen wording varies by account state, surface, and region and can change at any time. This is operational guidance, not legal advice. Compiled & annotated by CherryRed PR — cherryredpr.com.