

What is a model distillation attack?

A model distillation attack steals a model's capability without stealing its weights. The attacker sends large volumes of inputs to a target model through its API, collects the outputs, and trains a smaller student model to imitate them. The result is a cheaper copy that approximates the original, which is intellectual-property theft and usually a terms-of-service violation. It is a specific form of model extraction. Defenses include query rate limits and quotas, exposing less in each response, watermarking or perturbing outputs, detecting systematic harvesting, and enforcing terms of use.

HOW IT WORKS

01 How it works

The recipe is simple and scales with budget:

- Query at volume: send many diverse or carefully chosen prompts to the target model.
- Harvest the outputs, and if the API exposes them, the probabilities or token scores behind each answer.
- Train the student on those input and output pairs to match the target.

The more an API exposes, full text, long outputs, raw probabilities, the cheaper and closer the clone becomes.

02 Why it matters

A successful distillation costs the owner real value:

- Intellectual-property loss: the investment in building the model is approximated for a fraction of the cost.
- Terms-of-service violation: most providers forbid using outputs to train a competing model.
- Downstream risk: a local copy can be probed offline to find weaknesses in the original, or used to build attacks against it without rate limits.

HOW TO DEFEND

- Rate-limit and quota per authenticated identity, and require authentication in the first place.
- Expose less: avoid returning raw logits or token probabilities, and cap output detail where the product allows.
- Watermark or perturb outputs so a distilled copy carries a detectable trace.
- Detect systematic harvesting: high-volume, unusually diverse, automated query patterns from one account.
- Enforce terms of use against accounts that show extraction behaviour.

SOURCES

- [1] Adversarial Threat Landscape for AI Systems
- [2] AI Risk Management Framework
- [3] OWASP Top 10 for LLM Applications

Find out how much of your model an outsider could clone.

securelayer7.net/learn/ai-security/model-distillation-attack

[Open online](https://securelayer7.net/learn/ai-security/model-distillation-attack)