

THE INFRINGING MACHINE

Commercial Exploitation and AI Copyright Liability

Travis Gilly¹

Real Safety AI Foundation

Preprint (working draft) — April 2026

Target Venue: SSRN Preprint / Subsequent law review submission

1.

Travis Gilly is Founder and Executive Director of the Real Safety AI Foundation. Correspondence: t.gilly@ai-literacy-labs.org.

ABSTRACT

When AI systems generate outputs that reproduce copyrighted works, the liability analysis is more straightforward than prevailing discourse suggests. AI companies know their models memorize copyrighted content. They sell access to those models. Commercial exploitation of copyrighted content known to exist within a product constitutes infringement under existing doctrine.

This Article cuts through the epistemological debates surrounding training data, neural network “learning,” and machine cognition. These questions, while intellectually interesting, are largely irrelevant to the liability determination. What matters is that companies possess the right and ability to supervise infringing output, profit directly from that output, and have chosen not to prevent it effectively. Under established precedent, this constitutes vicarious and contributory infringement.

The Article makes three contributions. First, it reframes the liability question around commercial exploitation rather than training data provenance. Second, it dismantles the *Sony* defense by demonstrating that AI services are fundamentally different from neutral consumer products: they contain copyrighted content embedded in their parameters, and they operate as ongoing services with real-time control rather than products sold at arms-length. Third, it synthesizes existing precedent from international copyright rulings, vicarious liability doctrine, and contributory infringement case law to establish a unified analytical approach to AI copyright enforcement.

Keywords: AI copyright, vicarious liability, contributory infringement, inducement, *Grokster*, DMCA safe harbor, *Sony* doctrine, memorization, training data, secondary liability, *GEMA v. OpenAI*, generative AI, commercial exploitation, machine-generated works, Stanford extraction study.

CONTENTS

I	Introduction: The Liability Case in Brief	3
A	The Core Proposition	3
B	Why the Training Data Debate Is a Distraction	3
C	The Contribution of This Article	4
D	Roadmap	4
II	The Liability Framework: Vicarious and Contributory Infringement	5
A	Vicarious Liability	5
1	Right and Ability to Supervise	5
2	Direct Financial Interest	6
3	Application	7
B	Contributory Infringement	7
1	Knowledge	7
2	Material Contribution	8
C	Inducement Liability	8
D	The “Volition” Question	9
III	International Precedent: GEMA v. OpenAI	10
A	The German Ruling	10
B	Alignment with U.S. Vicarious Liability Doctrine	11
C	Applicability to U.S. Doctrine	11
IV	The Sony Defense and Why It Fails	12
A	The Defense Stated	12
B	Ground One: Pre-Loaded Content	12
C	Ground Two: Services Versus Products	13
D	The Combined Effect	14
V	Additional Defenses and Responses	14
A	The Guardrail Defense	14
B	The “Just a Tool” Defense	15
C	The Transformative Use Defense	15
D	The “Probabilistic System” Defense	16
E	The DMCA Safe Harbor Defense	17
VI	Remedies	18
A	Injunctive Relief	18
B	Damages	19
C	Structural Remedies	19
VII	Conclusion	20
	Acknowledgements	21

I. INTRODUCTION: THE LIABILITY CASE IN BRIEF

A. *The Core Proposition*

Stanford researchers recently demonstrated that commercial AI models can reproduce copyrighted books nearly verbatim.² Claude outputs 95.8% of *Harry Potter and the Sorcerer’s Stone*. Gemini reproduces 76.8%. These are not approximations or paraphrases; they are near-verbatim extractions of copyrighted text.

The discourse that follows such findings tends toward epistemological complexity. Commentators will argue over whether web scraping constitutes the same act as direct file ingestion, what “memorization” means for a neural network, and whether probabilistic reconstruction qualifies as “copying” in a legally meaningful sense.

This Article argues that these debates, while intellectually productive, are largely immaterial to the liability determination. The operative facts are these: AI companies know their models memorize copyrighted content (the research proves it, and their own remediation efforts confirm it); they charge money for access to those models (subscription fees, API pricing, and enterprise contracts); and they possess the technical capability to prevent infringing outputs (demonstrated by the content moderation systems they already operate for other categories of harmful content). The method by which copyrighted content entered the model is a question of factual interest. The commercial sale of a system known to reproduce that content is a question of legal consequence.

B. *Why the Training Data Debate Is a Distraction*

The current discourse focuses on a set of questions that, while important for understanding AI systems, do not determine liability. Did companies train directly on book files, or did

2.

Ahmed Ahmed, A. Feder Cooper, Sanmi Koyejo & Percy Liang, *Extracting Books from Production Language Models*, arXiv:2601.02671 (Jan. 6, 2026), available at https://papers.ssrn.com/sol3/papers.cfm?abstract_id=6050534.

models learn from web fragments? Does the methodology prove intentional ingestion? What constitutes “copying” versus “learning”? Is memorization equivalent to storage?

Consider three scenarios. In the first, a company feeds a copyrighted PDF directly into its training pipeline; the model memorizes it; the company sells access. In the second, a company trains on the public internet; the model learns to reconstruct copyrighted text from billions of scattered fragments; the company sells access. In the third, a model aggregates copyrighted content from user interactions over time; the company sells access.

In all three scenarios, the company sells access to a system that outputs copyrighted content on demand. The mechanism of memorization varies. The commercial exploitation does not.

C. The Contribution of This Article

This Article makes three contributions to AI copyright liability doctrine.

First, it reframes the liability question around commercial exploitation rather than training data provenance, cutting through the epistemological complexity that currently advantages defendants.

Second, it demonstrates that the *Sony* defense fails for AI systems on two independent grounds: AI services contain copyrighted content already embedded in their parameters (unlike neutral devices), and they operate as ongoing services with real-time control (unlike products sold at arms-length).

Third, it synthesizes existing legal doctrine, including vicarious liability, contributory infringement, and international copyright rulings, into a unified analytical approach to AI copyright liability.

D. Roadmap

Part II establishes the affirmative liability case based on vicarious and contributory infringement doctrine. Part III addresses international precedent demonstrating that courts have

already begun applying these principles to AI systems. Part IV dismantles the *Sony* defense. Part V addresses additional defenses and explains why each fails. Part VI examines remedies. Part VII concludes with implications for AI governance.

II. THE LIABILITY FRAMEWORK: VICARIOUS AND CONTRIBUTORY INFRINGEMENT

A. Vicarious Liability

Vicarious liability for copyright infringement attaches when a defendant has (1) the right and ability to supervise the infringing activity and (2) a direct financial interest in it.³ In *Shapiro*, the Second Circuit held that when the right and ability to supervise coalesce with an obvious and direct financial interest in the exploitation of copyrighted materials, liability attaches even absent actual knowledge that the copyright monopoly is being impaired.⁴ AI companies satisfy both prongs.

1. Right and Ability to Supervise

AI companies demonstrably possess the right and ability to supervise model outputs. The evidence is their own conduct.

When researchers expose memorization vulnerabilities, companies update their models. OpenAI, Anthropic, Google, and others have implemented content filtering systems (commonly called “guardrails”) designed to prevent the generation of violent content, hate speech, malware, and child sexual abuse material.⁵ Content policies prohibit various categories of

3.

Shapiro, Bernstein & Co. v. H.L. Green Co., 316 F.2d 304, 307 (2d Cir. 1963).

4.

Id. at 307.

5.

These systems include input filters that screen prompts, output filters that evaluate generated text before delivery, and model-level alignment techniques such as reinforcement learning from human feedback (RLHF). Their existence demonstrates technical capacity; their scope demonstrates deliberate choice about which categories of harmful output to prevent.

generation. Safety filters intercept requests in real time. Models are updated to refuse certain prompts.

By implementing these controls for other categories of content, companies establish that they possess the technical infrastructure to supervise and control output. The question is not capability; it is whether they have exercised that capability adequately with respect to copyright.

They have not. The Stanford extraction study tested four major production models and found that two of four did not even attempt to prevent verbatim reproduction of copyrighted books. The remaining two could be bypassed with known adversarial techniques.⁶ If companies can deploy real-time content filtering for violence, hate speech, and malware, they possess the technical capacity to implement analogous systems for verbatim reproduction of copyrighted text. That they have not done so reflects a business decision, not a technical limitation.

2. Direct Financial Interest

The financial interest prong is satisfied on its face. AI companies charge subscription fees (\$20/month for ChatGPT Plus, \$20/month for Claude Pro), API access fees (per-token pricing), and enterprise contracts (often valued in the millions of dollars annually).

The infringing capability is not incidental to the commercial product; it is inseparable from it. Users pay for access to models that generate text on demand. When that generated text reproduces copyrighted works, the reproduction occurs within the commercial service the user has paid to access. The company derives revenue from the very system that produces infringing output.

6. Ahmed et al., *supra* note 2. Gemini 2.5 Pro and Grok 3 required no jailbreaking to extract copyrighted text. Claude 3.7 Sonnet and GPT-4.1 exhibited refusal mechanisms, but these could be circumvented using a Best-of-N jailbreak technique.

3. Application

Both prongs are satisfied. AI companies possess the right and ability to supervise infringing outputs, as demonstrated by the existence of content controls for other categories of harmful output. AI companies derive direct financial benefit from access to models that generate those outputs, as demonstrated by their commercial pricing.

This analysis follows directly from *Shapiro, Bernstein & Co. v. H.L. Green Co.*,⁷ in which the Second Circuit imposed vicarious liability on a department store for the infringing sales of its concessionaire, and *A&M Records, Inc. v. Napster, Inc.*,⁸ in which the Ninth Circuit held that Napster’s ability to block users and police its system, combined with its financial interest in the infringing activity, established vicarious liability.

B. Contributory Infringement

Contributory infringement occurs when a defendant (1) knows of the infringing activity and (2) induces, causes, or materially contributes to it.⁹

1. Knowledge

AI companies possess actual knowledge that their models can reproduce copyrighted content. This element is no longer subject to genuine dispute.

Researchers have published peer-reviewed studies documenting memorization in large language models.¹⁰ Plaintiffs have filed lawsuits with exhibits showing verbatim reproduction. Most significantly, the companies themselves have acknowledged the phenomenon: Ope-

7.

316 F.2d 304 (2d Cir. 1963).

8.

239 F.3d 1004 (9th Cir. 2001).

9.

Gershwin Publ’g Corp. v. Columbia Artists Mgmt., Inc., 443 F.2d 1159, 1162 (2d Cir. 1971).

10.

In addition to Ahmed et al., *supra* note 2, see Nicholas Carlini et al., *Extracting Training Data from Large Language Models*, 30th USENIX Security Symposium (2021); A. Feder Cooper et al., *Extracting Memorized Pieces of (Copyrighted) Books from Open-Weight Language Models*, arXiv:2505.12546 (2025). Similar memorization and reproduction of copyrighted visual content has been documented in image generation models, extending the analysis presented here beyond text.

nAI’s public response to the *New York Times* litigation characterized instances of verbatim reproduction as “regurgitation” bugs the company was actively working to eliminate.¹¹

A party cannot work to eliminate a phenomenon it does not know exists. The remediation framing is itself an admission of knowledge.

2. *Material Contribution*

Companies do not merely know about infringement; they provide the complete infrastructure that makes it possible. They train the models that memorize the content. They host those models on their servers. They provide the API endpoints and user interfaces through which users access the models. They process the payments. They deliver the infringing output to the user’s screen.

Without the company’s active and ongoing participation at every stage, the infringement cannot occur. This satisfies the material contribution element.

C. *Inducement Liability*

Beyond vicarious and contributory liability, the Supreme Court has recognized a third form of secondary liability: inducement. In *MGM Studios, Inc. v. Grokster, Ltd.*, the Court held that a party who distributes a product with the purpose of promoting its infringing use, as shown by affirmative acts designed to foster infringement, is liable for the resulting infringement by third parties.¹²

Grokster matters here for two reasons. First, it confirms that the *Sony* safe harbor does not shield a distributor who knows of infringement and fails to act. The Court held that *Sony*’s rule “was never meant to foreclose rules of fault-based liability derived from the common law.”¹³

11.

OpenAI, *OpenAI and Journalism* (Jan. 8, 2024), <https://openai.com/index/openai-and-journalism/>.

12.

MGM Studios, Inc. v. Grokster, Ltd., 545 U.S. 913, 936–37 (2005).

13.

Id. at 934–35.

Second, the *Grokster* Court identified indicia of inducement relevant here: targeting known infringers, failing to develop filtering tools, and a business model dependent on infringing use.¹⁴ Each has an analog in AI deployment. Companies know their models reproduce copyrighted content; they have not deployed filtering tools adequate to prevent verbatim reproduction despite possessing the capacity to do so for other content categories; and the commercial value of the service derives substantially from generative capabilities that include infringing output.

The *Grokster* standard complements rather than replaces the vicarious and contributory analyses. A finding of inducement reinforces the conclusion that *Sony*’s safe harbor does not shield AI companies from liability for foreseeable infringing outputs.

D. The “Volition” Question

Courts have sometimes required “volitional conduct” for direct infringement, holding that purely automated processes do not satisfy this requirement.¹⁵ In *CoStar*, the Fourth Circuit concluded that an internet service provider that passively hosted user-uploaded content lacked the volitional conduct necessary for direct infringement liability.¹⁶ The Second Circuit extended this reasoning in *Cartoon Network (Cablevision)*, holding that a remote DVR service lacked the volitional conduct required for direct infringement because the copies at issue were created at the direction of the subscribers who initiated the recordings.¹⁷ AI companies may invoke this line of reasoning to argue that model outputs are automated responses to user prompts rather than volitional company conduct.

This argument fails for three reasons.

14.

Id. at 939–40.

15.

CoStar Grp., Inc. v. LoopNet, Inc., 373 F.3d 544 (4th Cir. 2004); *Cartoon Network LP, LLLP v. CSC Holdings, Inc.*, 536 F.3d 121 (2d Cir. 2008).

16.

CoStar, 373 F.3d at 550–51.

17.

Cartoon Network, 536 F.3d at 131. The Ninth Circuit has applied similar reasoning to automated indexing technologies. See *Perfect 10, Inc. v. Amazon.com, Inc.*, 508 F.3d 1146, 1160–61 (9th Cir. 2007) (concluding that Google’s automated image indexing did not constitute direct infringement).

First, neither vicarious nor contributory liability requires the defendant to personally perform the infringing act. The defendant need only supervise and fail to prevent (vicarious) or materially contribute to and fail to stop (contributory) infringement by others. The volition cases addressed *direct* infringement liability only; they do not immunize defendants from the secondary liability theories advanced here.

Second, the volition doctrine assumes that the service provider operates a neutral system in which users, not the provider, select the content to be reproduced. The Cablevision subscriber chose which broadcast to record from programming he could already access. The CoStar user chose which image to upload. AI systems are categorically different: the copyrighted content is already present in the model’s parameters before any user interaction, incorporated through training decisions the company made.

Third, AI companies exercise volition at every meaningful decision point in the chain: they chose which copyrighted works to include in training data, they chose to deploy commercially, they chose not to implement effective reproduction filters, and they continue to choose to operate with knowledge that their systems produce infringing output. These are volitional business decisions, not the passive automation at issue in *CoStar* or *Cablevision*.

III. INTERNATIONAL PRECEDENT: GEMA V. OPENAI

A. *The German Ruling*

In November 2025, the Munich I Regional Court ruled in *GEMA v. OpenAI* that memorization of protected song lyrics within ChatGPT’s model parameters constitutes copyright infringement under German and EU law.¹⁸

The court assigned liability to OpenAI as the commercial operator, reasoning that decisions regarding model architecture and dataset selection placed responsibility with the

¹⁸.

GEMA v. OpenAI, Munich I Regional Court, Case No. 42 O 14139/24 (Nov. 11, 2025). The court applied §§ 15, 16, and 19a of the German Copyright Act (*Urheberrechtsgesetz*) and Articles 2 and 3 of the EU InfoSoc Directive 2001/29/EC.

provider rather than with end-users.¹⁹ The court further held that the text-and-data-mining exception under Article 4 of the EU DSM Directive did not extend to memorization or reproduction of entire works, finding that such uses exceed the analytical purpose the exception was designed to serve.

B. Alignment with U.S. Vicarious Liability Doctrine

The German court’s reasoning maps onto the vicarious liability analysis with notable precision:

1. OpenAI made architectural decisions that determined what content the model would memorize (establishing supervision capacity);
2. OpenAI selected the training datasets (establishing knowledge of the content’s presence);
3. OpenAI chose to make the model commercially available with memorized content intact (establishing financial interest); and
4. Therefore, OpenAI bears liability as the commercial operator.

This reasoning is output-focused. The court did not require the plaintiff to demonstrate precisely how memorization occurred at a technical level or to identify which specific training files were ingested. It examined the output (recognizable copyrighted lyrics), the commercial context (a paid service), and the operator’s control (architecture and deployment decisions), and found these sufficient to establish liability.

C. Applicability to U.S. Doctrine

German copyright law operates under different statutory frameworks than U.S. law. *GEMA* is not binding precedent in American courts.

¹⁹.

Id. The court applied provider liability principles (*Störerhaftung*) and rejected OpenAI’s defense that outputs were user-initiated and therefore beyond the company’s control.

However, *GEMA* demonstrates that existing liability doctrines can be applied to AI systems without requiring novel legal theories. The core reasoning, that commercial operators bear responsibility for the outputs of systems they design, deploy, and monetize, is equally applicable under U.S. vicarious and contributory liability doctrine.

GEMA serves as persuasive authority indicating the direction in which courts are moving internationally. U.S. courts confronting similar facts should reach similar conclusions under existing domestic doctrine.

IV. THE SONY DEFENSE AND WHY IT FAILS

A. *The Defense Stated*

AI companies will invoke *Sony Corp. of America v. Universal City Studios, Inc.*²⁰ The Supreme Court held that selling a device capable of copyright infringement is permissible if the device is “capable of substantial noninfringing uses.”²¹ The VCR could record copyrighted television programs, but it could also record home videos, time-shift authorized content, and serve other legitimate purposes.

ChatGPT can draft emails, assist with programming, answer factual questions, and generate original content. These are substantial noninfringing uses. The *Sony* defense, the argument goes, should therefore apply.

This defense fails for AI systems on two independent grounds.

B. *Ground One: Pre-Loaded Content*

A VCR ships as an empty device. The user provides whatever content is subsequently recorded. The manufacturer sells a neutral tool; any infringement that follows is the user’s act performed with the user’s materials.

²⁰.
464 U.S. 417 (1984).

²¹.
Id. at 442.

An AI model does not ship empty. It contains copyrighted content already compressed into its weight parameters. The memorization has occurred before the user sends a single prompt. When a user subscribes to Claude or ChatGPT, the copyrighted content is already present within the system they are accessing.

The distinction is dispositive. *Sony* protects manufacturers of neutral tools. AI models are not neutral tools; they are systems with copyrighted content already embedded at the point of commercial deployment. The user does not supply the copyrighted material. The company does, incorporated into the product before any user interaction occurs.

Sony's rationale, that manufacturers cannot control what users do with neutral devices after sale, does not apply when the manufacturer has already incorporated infringing content into the device before sale.

C. Ground Two: Services Versus Products

Sony addressed products sold at arms-length. The doctrinal framework assumed a discrete transaction: manufacturer sells device, consumer takes possession, manufacturer has no ongoing relationship with use. Sony could not monitor what consumers did in their living rooms with their VCRs.

AI systems operate as services, not products. Every prompt flows through the company's servers. Every output is generated on company infrastructure. The company maintains real-time control over every interaction.

AI service providers can filter outputs in real time, refuse specific requests, update model behavior through over-the-air changes, monitor usage patterns, terminate access for policy violations, and implement new content moderation systems at any point. They exercise these capabilities daily for content categories they have prioritized: violence, hate speech, malware, and child sexual abuse material.

Sony protected manufacturers precisely because they lacked the practical ability to prevent downstream infringement. AI service providers demonstrably possess that ability. They

exercise it continuously for other categories of harmful content. They have not exercised it effectively for copyright.

A company cannot simultaneously claim *Sony* protection on the ground that it lacks control over user activity and market sophisticated content moderation capabilities that demonstrate precisely such control. These positions are irreconcilable. If the company possesses the ability to control outputs, the predicate for *Sony* protection does not exist.

D. The Combined Effect

Either ground independently defeats the *Sony* defense. Together, they are conclusive.

AI services contain copyrighted content already embedded in their parameters, unlike the neutral devices at issue in *Sony*. AI services operate with ongoing real-time control over every output, unlike products sold at arms-length. *Sony*'s rationale, protecting manufacturers who cannot control downstream use of neutral devices, is simply inapposite.

V. ADDITIONAL DEFENSES AND RESPONSES

A. The Guardrail Defense

AI companies may argue that they have implemented safety measures designed to prevent reproduction of copyrighted content.

The empirical evidence does not support this defense. The Stanford extraction study found that two of four major production models had no mechanisms whatsoever to prevent verbatim reproduction of copyrighted books. The remaining two had mechanisms that could be circumvented through known adversarial techniques.²²

More fundamentally, the existence of guardrails cuts against this defense rather than supporting it. By implementing content filtering systems, companies establish that they possess the technical capability to control outputs, satisfying the supervision prong of vicarious lia-

²².

Ahmed et al., *supra* note 2.

bility. The demonstrated inadequacy of those systems with respect to copyright establishes that the company has not discharged the responsibility that supervision capacity creates.

B. The “Just a Tool” Defense

Companies may characterize their AI systems as sophisticated tools, analogous to word processors or search engines, and argue that tools do not bear liability for user conduct.

If AI systems are tools, then the *Sony* analysis applies, and it fails for the reasons stated above.

The tool characterization also fails on its own terms. A conventional tool executes a user’s instructions. An AI model generates outputs based on probabilistic processes that incorporate content the user did not supply and may not have intended to elicit. When a user requests “a story about a boy wizard at a school for magic,” the model, not the user, selects the specific text to generate. If that text reproduces *Harry Potter*, the user did not instruct the model to infringe; the model selected an infringing output path from among many alternatives.

This creates a dilemma for the company. If it characterizes the model as a tool, it accepts responsibility for what the tool produces. If it characterizes the output as the model’s own generation, it acknowledges a degree of operational autonomy that does not eliminate vicarious liability for outputs the company profits from and possesses the ability to supervise.

C. The Transformative Use Defense

Companies argue that AI training and generation are inherently transformative, that models learn abstract patterns rather than copy specific content, and that outputs constitute new creative works.

This argument has force for certain categories of AI output. A model that synthesizes information from multiple sources to answer a factual question may produce transformative

output. A model that generates visual content in styles not attributable to any single artist may produce transformative output.

But transformativeness is evaluated on an output-by-output basis, not system-wide.²³ When ChatGPT reproduces paragraphs of a *New York Times* article verbatim, that specific output is not transformative.²⁴ When Claude outputs 95.8% of *Harry Potter and the Sorcerer’s Stone*, that specific output is not transformative.

Fair use analysis under 17 U.S.C. § 107 examines the purpose and character of the use, the nature of the copyrighted work, the amount and substantiality of the portion used, and the effect on the potential market for the original.²⁵ Verbatim reproduction of creative works through a commercial service, competing directly with the market for the originals, fails all four factors.

The existence of some transformative outputs does not immunize infringing outputs. Each output stands or falls on its own merits.

D. The “Probabilistic System” Defense

Companies may argue that AI outputs are probabilistic, that the specific content generated in response to any given prompt is unpredictable, and that liability cannot attach to outcomes the company could not have foreseen.

Unpredictability of specific instances does not eliminate liability for foreseeable categories of harm. A trucking company cannot escape liability for accidents by arguing that the time and location of any specific collision are unpredictable. A pharmaceutical manufacturer does not escape product liability because individual adverse reactions cannot be predicted with certainty.²⁶

23.

See *Campbell v. Acuff-Rose Music, Inc.*, 510 U.S. 569, 581 (1994) (holding that the inquiry “may be guided by the examples given in the preamble to § 107” and must be applied to the specific use at issue).

24.

Complaint, *N.Y. Times Co. v. Microsoft Corp.*, No. 1:23-cv-11195 (S.D.N.Y. Dec. 27, 2023), Exs. J, K (documenting near-verbatim reproduction of Times articles).

25.

17 U.S.C. § 107.

26.

Liability attaches to the enterprise that creates foreseeable categories of harm, not to the prediction of specific instances within those categories.

AI companies know their models can reproduce copyrighted content. They know this because researchers have demonstrated it empirically, because plaintiffs have documented it in litigation, and because they have invested resources in attempting to remediate it. The specific infringing output may be unpredictable in any given interaction; the category of infringement is entirely foreseeable.

E. The DMCA Safe Harbor Defense

AI companies may seek refuge in the safe harbor provisions of 17 U.S.C. § 512, which shield qualifying online service providers from monetary liability for copyright infringement under specified conditions.²⁷ This defense fails at the threshold.

Section 512(c), the provision most plausibly applicable, protects service providers from liability for infringing material “residing on a system or network controlled or operated by or for the service provider” that was placed there “at the direction of a user.”²⁸ Memorized training data is not material residing on the system at the direction of a user. It is material incorporated into the model’s parameters as a consequence of decisions the company made regarding training data selection and model architecture. The user does not upload *Harry Potter* to the model; the company did, by training on content it knew was copyrighted, and compressed the resulting reproduction capability into the commercial product.

The remaining safe harbors are similarly inapposite. Section 512(a) applies to transitory digital network communications; AI outputs are generated by the provider’s own system rather than transmitted between third parties. Section 512(b) applies to intermediate and

See, e.g., *Greenman v. Yuba Power Prods., Inc.*, 59 Cal.2d 57 (1963) (establishing strict product liability for foreseeable categories of harm regardless of the manufacturer’s ability to predict specific instances).

²⁷.

17 U.S.C. § 512. The four distinct safe harbors address transitory digital network communications (§ 512(a)), system caching (§ 512(b)), information residing on systems or networks at the direction of users (§ 512(c)), and information location tools (§ 512(d)).

²⁸.

17 U.S.C. § 512(c)(1).

temporary storage of material transmitted by others; AI-generated text is not cached transmission but generation from memorized parameters. Section 512(d) applies to information location tools that link users to material residing on other networks; AI models do not function as search indices or directories.

Even if a court were inclined to apply Section 512(c) by analogy, the statute conditions its protection on the absence of actual knowledge of infringing activity, the absence of “red flag” awareness of facts making infringement apparent, and expeditious removal upon notification.²⁹ AI companies possess actual knowledge that their models reproduce copyrighted content, as established in Part II.B above. The red flag standard articulated in *Viacom International, Inc. v. YouTube, Inc.*³⁰ is satisfied by the published extraction studies, the litigation record, and the companies’ own public acknowledgments of the memorization phenomenon.

The DMCA was enacted to protect neutral intermediaries whose systems carry user-generated content. It was not designed to, and does not, shield commercial providers whose systems themselves contain and reproduce copyrighted expression as a consequence of deliberate architectural choices.

VI. REMEDIES

A. Injunctive Relief

Injunctive relief is the most probable remedy in practice. Courts have broad authority to grant injunctions in copyright cases,³¹ and have exercised that authority against technology platforms in analogous contexts.³² If a model is found to infringe, courts could enjoin

²⁹.

¹⁷ U.S.C. § 512(c)(1)(A)–(C).

³⁰.

Viacom Int’l, Inc. v. YouTube, Inc., 676 F.3d 19 (2d Cir. 2012).

³¹.

¹⁷ U.S.C. § 502(a) (“Any court having jurisdiction of a civil action arising under this title...may...grant temporary and permanent injunctions on such terms as it may deem reasonable to prevent or restrain infringement of a copyright.”).

³².

commercial deployment until infringement is cured, require licensing of underlying works as a condition of continued operation, mandate output filtering to prevent reproduction of specific protected works, or order regular auditing and reporting on memorization.

These remedies carry substantial practical force. Commercial AI deployment generates significant revenue. The prospect of injunction pending licensing creates powerful incentives for pre-litigation agreements with rights holders.

B. Damages

Damages in copyright cases include actual damages, the defendant’s profits attributable to the infringement, or statutory damages at the plaintiff’s election.³³

Actual damages might include lost licensing fees that would have been paid for authorized commercial use of the copyrighted works.

The defendant’s profits might include revenue attributable to infringing capabilities, calculated as a reasonable royalty for unauthorized commercial exploitation of copyrighted content.

Statutory damages provide a floor regardless of the plaintiff’s ability to prove actual harm, and can be substantial where infringement is found to be willful.³⁴

C. Structural Remedies

For systematic infringement, courts may order structural remedies: training data audits to identify the presence of copyrighted content, output monitoring systems to detect and prevent reproduction, licensing mandates requiring agreements with rights-holder organizations, and periodic compliance reporting.

See, e.g., *A&M Records, Inc. v. Napster, Inc.*, 239 F.3d 1004, 1027 (9th Cir. 2001) (affirming preliminary injunction requiring Napster to prevent users from sharing copyrighted files).

³³.

17 U.S.C. § 504.

³⁴.

17 U.S.C. § 504(c)(2) (providing for statutory damages of up to \$150,000 per work infringed in cases of willful infringement).

These remedies address the ongoing nature of AI operations. Unlike a discrete act of infringement (a single unauthorized publication, for instance), AI models continue generating outputs indefinitely across millions of user interactions. Structural remedies ensure continuing compliance with the court’s determination.

VII. CONCLUSION

The liability case against AI companies for copyright infringement is more straightforward than prevailing discourse suggests.

AI companies know their models memorize copyrighted content. The empirical research establishes this. Their own remediation efforts confirm it. The knowledge element is satisfied.

They sell access to those models through subscription fees, API pricing, and enterprise contracts. The financial interest element is satisfied.

They possess the technical capability to supervise and control model outputs, as demonstrated by the content moderation systems they operate for other categories of harmful content. The supervision element is satisfied.

They provide the complete infrastructure, from training to hosting to delivery, without which the infringement cannot occur. The material contribution element is satisfied.

This is vicarious, contributory, and induced infringement under the doctrine established in *Shapiro*,³⁵ *Gershwin*,³⁶ *Napster*,³⁷ and *Grokster*.³⁸

The *Sony* defense fails. AI services contain copyrighted content embedded in their parameters, unlike neutral devices. AI services operate with ongoing real-time control over every output, unlike products sold at arms-length. *Sony*’s rationale does not apply, and

³⁵.

Shapiro, Bernstein & Co. v. H.L. Green Co., 316 F.2d 304 (2d Cir. 1963).

³⁶.

Gershwin Publ’g Corp. v. Columbia Artists Mgmt., Inc., 443 F.2d 1159 (2d Cir. 1971).

³⁷.

A&M Records, Inc. v. Napster, Inc., 239 F.3d 1004 (9th Cir. 2001).

³⁸.

MGM Studios, Inc. v. Grokster, Ltd., 545 U.S. 913 (2005).

Grokster forecloses any attempt to stretch *Sony* beyond its limits. The DMCA safe harbor is unavailable because memorized training data is not content stored at the direction of a user.

The epistemological debates about training data provenance, neural network cognition, and machine “learning” invite complexity where the doctrine requires none. The liability determination rests on established elements: knowledge, financial interest, supervisory capacity, and material contribution. Each element is satisfied on the present facts.

ACKNOWLEDGEMENTS

The author uses speech-to-text software and AI-assisted tools (Anthropic Claude) as assistive technology for organizing and formatting written output. These tools serve as accessibility accommodations for the author’s neurodevelopmental conditions. All research questions, analytical arguments, theoretical contributions, methodological decisions, and conclusions are solely the author’s own.