

IP Alert | Biggest Public Copyright Settlement in History: Anthropic and Book Authors Announce Terms to Settle Lawsuit Over Pirated Material Used to Train LLM

By Victoria Webb and Kirk Sigmon

On September 5, 2025, artificial intelligence (“AI”) company Anthropic publicly agreed to pay at least \$1.5 billion for the infringing use of around 500,000 copyrighted works as part of training its Claude Large Language Model (“LLM”) products. The case, *Bartz et al. v. Anthropic PBC*, No. 24-cv-05417 (N.D. Cal.), was already of major legal significance because of the Court’s [summary judgment ruling earlier this year that use of copyrighted material, standing alone, for machine learning training was fair use, but use of pirated copies for machine learning training was not](#). Now, if this “landmark” class action settlement is approved, the case will also enter the history books as the largest public copyright settlement in history.

Background: Pirated Works Used to Train Claude LLMs

This class action hinges around Anthropic’s use of certain “shadow libraries” (pirated repositories of digital texts, including copyrighted material) for training their LLMs, such as their Claude LLM products. The pirated material used by Anthropic includes the Library Genesis (“LibGen”) and Pirate Library Mirror (“PiLiMi”) shadow libraries, which were apparently downloaded (via torrents) by Anthropic and used without authorization of any of the copyright owners of the works contained in those shadow libraries. Anthropic’s use of those pirated works [had already been found “inherently, irredeemably infringing” of authors’ copyrights by a Northern District of California judge in response to a summary judgment motion](#), kicking off a laborious process of multi-month-long settlement discussions between the parties.

According to the Motion for Preliminary Approval of Class Settlement filed September 5, 2025, the parties undertook significant discovery on the road to settlement—including exchanging “hundreds of thousands of pages of documents, litigating over one dozen discovery motions, conducting 20 depositions,” and more. Indeed, the settlement seems to have been struck days before the close of fact discovery and as the parties were “actively preparing” for Anthropic’s CEO to be deposed.

In addition, the road to settlement seems to have involved input from other big players in the copyright space, such as the Authors Guild and the Association of American Publishers. This is notable, as many copyright membership and trade associations for rightsholders are

involved in their own ongoing AI-related copyright litigation (e.g., Authors Guild v. OpenAI et al., No. 1:23-cv-8292 (S.D.N.Y.)).

Settlement Details: \$3k per Work, Anthropic to Destroy Torrented Shadow Libraries

The settlement boils down to three “principal terms”:

1. Anthropic will pay the class a minimum of \$1.5 billion, totaling around \$3,000 per class work;
2. Anthropic agrees to destroy their copies of works acquired from shadow libraries such as the LibGen and PiLiMi data sets; and
3. Anthropic is only released for past training behavior, meaning that it cannot use the copyrighted works for future training without legally acquiring the works, and meaning that Anthropic is not protected from suits involving output from its LLMs.

Critically, the \$1.5 billion figure is a minimum, and might be increased by \$3,000 per work if additional works are identified and added. That said, the total funds might be reduced by, for example, settlement administration expenses, fee awards to the attorneys that worked on the case, and to other individuals (experts, special masters, industry experts) hired as part of the case. Pragmatically speaking, this may mean that authors receive substantially less than the full \$3,000 per work, especially if there are multiple authors per work (as they'd end up essentially splitting the award). And, in any event, the \$3,000 figure could be substantially less than might have been available via statutory damages available under 17 U.S.C. § 504(c)(1) (“with respect to any one work...a sum of not less than \$750 or more than \$30,000 as the court considers just”) or 17 U.S.C. § 504(c)(2) (if the Court finds willful infringement, “the court in its discretion may increase the award of statutory damages to a sum of not more than \$150,000”).

Settlement Requires Marketing Online, via Facebook, Instagram, Reddit

Recognizing that it might be hard to locate the copyright owners of all of the works in the LibGen and PiLiMi datasets (after all, most books don't exactly contain author telephone numbers or e-mail addresses), Anthropic is required to use a variety of strategies to help locate potential authors. In addition to opening a website where owners can file claims (www.AnthropicCopyrightSettlement.com), Anthropic has agreed to some particular marketing terms, including:

- Anthropic agrees to pay for Google ads targeting “Adults twenty-five (25) years of age or older” who are “in-market” for services such as book promotion services and/or literary agents, who have browsed websites such as selfpublishing.com or scribophile.com, who have searched Google for terms such as “manuscript” “amazon kdp” “Reedsy,” “critique circle,” “ASIN,” “indie author,” “Scribophile,” “author royalties,” “how to publish a book,” “book publishing,” and various AI-related terms.
- Anthropic also agrees to publish notice on Facebook and Instagram, directing such ads to individuals with job titles such as “Journalist,” “Writer,” “Online Publisher,” and the like.
- Anthropic also agrees to help advertise the settlement on Reddit, such as towards users “who participate in communities including r/selfpublishing, r/publishing, r/selfpublish, r/writers, and r/writing.”
- Anthropic is required to publish notice in publications such as Publishers Weekly, The Atlantic, the Toronto Star, The Globe and Mail, and La Presse.

Other Opportunities for Training Using Copyrighted Works May Exist

Despite the language of the settlement providing release for only “past . . . conduct,” it is possible that Anthropic could still find ways to use copyrighted material to train their Claude LLM products without the express consent of the copyright owner. As detailed in our [article regarding the aforementioned summary judgment ruling in this case](#), at least two judges in the Northern District of California have ruled that use of legally-acquired material (e.g., purchased e-books, purchased and scanned physical books) to train machine

learning models may comprise fair use. Along those lines, Anthropic could, like other LLM companies have done, legally purchase entire physical and/or virtual libraries and process those libraries to generate their own training data without receiving explicit permission from the copyright owners.

Anthropic Still Liable for Infringing Output

This settlement explicitly does not release Anthropic for allegedly infringing outputs. In other words, if Anthropic's LLMs are capable of creating copyright-infringing work (e.g., complete replications of an author's work, such as reproducing the entirety of their book or paper), they might still be vulnerable to an entirely separate copyright lawsuit. This is precisely currently at issue in a related Central District of California case, *Disney Enterprises, Inc. v. Midjourney Inc.*, No. 2:25-CV-05275 (C.D. Cal.), where complainants such as Universal, Disney, DreamWorks, Marvel, Lucasfilm, and Twentieth Century Fox allege that Midjourney (and, for instance, Midjourney's video generation functionality) can be used to create copyright-infringing content (e.g., videos that use popular Disney or Marvel characters).

Takeaway: Vetting Training Data Remains Extremely Important

This settlement may strike fear into the hearts of LLM developers, particularly those that might have used pirated material in the past. Eye-popping quantities of data are required to properly train LLMs, and that exposes those training LLMs to equally eye-popping damages should even a fraction of those works be acquired through unlawful means. Moreover, these types of damages could apply to unlawful use of any kind of training data: books, movies, pictures, or the like. Now, more than ever, AI companies or others training machine learning models should carefully ensure that their training data is legally acquired (e.g., legally purchased, acquired, and/or otherwise licensed).

Posted: September 5, 2025