

Joseph R. Saveri (State Bar No. 130064)
JOSEPH SAVERI LAW FIRM, LLP.
601 California Street, Suite 1505
San Francisco, California 94108
Telephone: (415) 500-6800
Facsimile: (415) 395-9940
Email: jsaveri@saverilawfirm.com

Attorneys for Plaintiffs

**UNITED STATES DISTRICT COURT
NORTHERN DISTRICT OF CALIFORNIA
SAN JOSE DIVISION**

**SUSANA MARTINEZ-CONDE, STEPHEN L.
MACKNIK,**

Plaintiffs,

v.

APPLE INC.,

Defendant.

Case No.

CLASS ACTION COMPLAINT

CLASS ACTION

DEMAND FOR JURY TRIAL

CONTENTS

I.	INTRODUCTION	1
II.	JURISDICTION AND VENUE	3
III.	DIVISIONAL ASSIGNMENT	4
IV.	PARTIES.....	4
A.	Plaintiffs	4
B.	Defendant.....	5
C.	Agents and Co-Conspirators	6
V.	FACTUAL ALLEGATIONS	6
A.	How large language models work.....	7
B.	Apple trained its OpenELM models on copyrighted works.	9
C.	Apple trained its Foundation Language Models on copyrighted works.	13
D.	Apple’s conduct impairs the market for Plaintiffs’ and Class members’ works...	18
VI.	CLASS ACTION ALLEGATIONS	19
VII.	CLAIMS	22
	COUNT ONE	22
VIII.	PRAYER FOR RELIEF	23
IX.	DEMAND FOR JURY TRIAL	24

1 Plaintiffs Susana Martinez-Conde and Stephen L. Macknik (“Plaintiffs”), on behalf of
2 themselves and all others similarly situated (the “Class,” as defined below), bring this class-action
3 complaint (“Complaint”) against Defendant Apple Inc. (“Apple” or “Defendant”).

4 I. INTRODUCTION

5 1. Apple infringed upon Plaintiffs’ and Class members’ copyrights by reproducing
6 their registered works without authorization as a part of amassing centralized databases of
7 training materials and using that data to train its “Apple Intelligence” AI models on Plaintiffs’ and
8 Class members’ copyrighted books and other works. Apple has created a set of generative AI
9 models collectively called Apple Intelligence that it provides to consumers in its phones, tablets
10 and personal computers. It used Plaintiffs’ and Class members’ copyrighted works without their
11 authorization, without compensating them to train and test their Apple Intelligence models, a
12 impermissible use far beyond any applicable license Apple has to sell such books to the users of its
13 products.

14 2. Apple Intelligence is comprised of multiple generative AI models. These include
15 but are not limited to “Apple Foundation Models,” a ~3 billion parameter on-device language
16 model, and a larger server-based language model, and Apple’s OpenELM models, a family of
17 “Open Efficient Language Models.”

18 3. Apple reproduced and used data sets that included Books3, a dataset of pirated,
19 copyrighted books that includes the published works of Plaintiffs and the Class, as training data
20 for its AI models. Apple used the books in Books3, among others, to train its OpenELM language
21 models and its Foundation Language Models.

22 4. Books3 is a notorious “shadow library,” a dataset of pirated, copyrighted books
23 that can be found in various places on the internet or shared and downloaded from pirate websites
24 and file sharing protocols like BitTorrent, the popular peer-to-peer (P2P) file sharing protocol for
25 copying and distributing infringing material. It is often found as a part of other datasets of pirated
26 books.

27 5. Along with knowingly using training datasets that include Books3 to train its
28 models, Apple uses “Applebot,” a web-crawling software program that copies mass quantities of

1 internet data (also known as “scraping”) to use as training data. Apple scraped data with
2 Applebot for nearly nine years before disclosing that it intended to use the scraped data to train its
3 AI systems. Web crawlers like Applebot scrape shadow libraries including but not limited to
4 Books3, that host millions of other unlicensed copyrighted books, including Plaintiffs’ and Class
5 members’ copyrighted works.

6 6. On information and belief, Apple also trained its models on unauthorized copies of
7 eBooks it sells to its users through Apple Books. Copying and using such eBook files for any
8 purpose beyond the explicit, limited scope of Apple’s license to sell them is copyright
9 infringement.

10 7. Generative AI models like those used in Apple Intelligence are only as good as the
11 training data on which they are trained. Bad writing in training data results in less valuable, less
12 useful AI models. Good writing in training data makes AI outputs better and models more
13 valuable. This is why Apple and other AI companies prioritize and use high quality writing, like
14 copyrighted works, to train and fine-tune their models.

15 8. Apple’s AI models were created and operate by first making unauthorized digital
16 copies of copyrighted works as a part of amassing central databases of enormous amounts of
17 textual works and images to use for various purposes, including as training data. Apple then
18 prepares such data for training AI models, which involves the creation of unauthorized copies and
19 the preparation of derivative works in training data sets. Apple trains a model by using computers
20 to analyze and record massive amounts of information about the relationships between words or
21 bits of words (“tokens”) in those works; using advanced mathematics and computer processing
22 to predict a sequence of words in response to a text prompt, based on that detailed information
23 about the relationship between tokens—the very creative expression found in those copyrighted
24 works; and fine-tuning and mid-training the model on the most desired texts with the best
25 writing, to achieve preferred model outcomes and outputs.

26 9. Plaintiffs and the Class are authors who have registered copyrights for their
27 published works. They did not consent to the unauthorized reproduction by Apple and use of
28 their works to be stored in databases for general use by Apple or for use in any Apple Intelligence

1 model, including the Foundation Intelligence Models and OpenELM language models. Such
2 pilfered intellectual property, along with being used for AI training and fine-tuning, is used for
3 testing model performance, and for the creation of filters to prevent model outputs containing
4 recited or regurgitated copyrighted materials from reaching the end user.

5 10. Plaintiffs and the Class did not consent to Apple making further reproductions of
6 their works or preparing altered derivative work based on their copyright works to use as training
7 data for Apple intelligence models. Apple prepared derivative works by processing and modifying
8 the raw training data they copied, including Plaintiffs' books, to create derivative training data
9 sets used in training their Apple Intelligence models.

10 11. The market for licensing AI training data is growing rapidly. Licensing deals to use
11 copyrighted works as training data between AI developers and publishers are regularly in the
12 news. Nevertheless, Apple did not compensate creators for use of their copyrighted works and
13 concealed the sources of their training datasets to evade legal scrutiny. Apple continues to retain
14 private AI training-data, including pirated books, to train its future models in various datasets
15 without seeking Plaintiffs' or Class members' consent or providing them compensation.

16 12. Apple has illegally copied Plaintiffs copyrighted works to train its AI models,
17 whose outputs compete with and dilute the market for those very works—works without which
18 Apple Intelligence would have far less commercial value. This conduct has damaged Plaintiffs'
19 and Class members' intellectual property. It deprived Plaintiffs and the Class of control over their
20 work, undermined the economic value of their copyrighted works, and positioned Apple to
21 achieve massive commercial success through unlawful means.

22 II. JURISDICTION AND VENUE

23 13. This Court has subject-matter jurisdiction under 28 U.S.C. §§ 1331, 1332(d),
24 1400(a) because this case arises under the Copyright Act (17 U.S.C. § 501).

25 14. Jurisdiction and venue are proper in this judicial district under 28 U.S.C. §§
26 1391(c)(2) and 1400(a) because Apple is headquartered in this district.

III. DIVISIONAL ASSIGNMENT

15. Pursuant to Civil Local Rules 3-2(c) and 3-2(e), assignment of this case to the San Jose Division is proper because Apple is located in Santa Clara County where a substantial part of the events giving rise to Plaintiffs' and Class members' claims occurred.

IV. PARTIES

A. Plaintiffs

16. **Plaintiff Susana Martinez-Conde** is a Professor of Ophthalmology, Neurology, and Physiology & Pharmacology at SUNY Downstate Health Sciences University. Professor Martinez-Conde received a BSc in Experimental Psychology from Universidad Complutense de Madrid and a Ph.D in Medicine and Surgery from the Universidade de Santiago de Compostela in Spain. She was a postdoctoral fellow with the Nobel Laureate Prof. David Hubel, and then an Instructor in Neurobiology, at Harvard Medical School. Professor Martinez-Conde writes frequently for *Scientific American* and previously had a regular column in *Scientific American: MIND* on the neuroscience of illusion. She is the 2014 recipient of the Science Educator Award, a prestigious prize given by the Society for Neuroscience (30,000 members) to an outstanding neuroscientist who has made significant contributions to educating the public.

17. Prof. Martinez-Conde's research has been featured in print in The New York Times, The New Yorker, The Wall Street Journal, Wired, The LA Chronicle, The Times (London), The Chicago Tribune, The Boston Globe, Der Spiegel, etc., and in radio and TV shows, including Discovery Channel's Head Games and Daily Planet shows, NOVA: scienceNow, CBS Sunday Morning, NPR's Science Friday, and PRI's The World. She works with international science museums, foundations and nonprofit organizations to promote neuroscience education and communication.

18. **Plaintiff Stephen L. Macknik** is a Professor of Ophthalmology, Neurology, and Physiology & Pharmacology at SUNY Downstate Health Sciences University. Professor Macknik is Professor Martinez-Conde's writing partner and husband. Together, they authored the international bestselling book Sleights of Mind: What the Neuroscience of Magic Reveals About Our Everyday Deceptions, which has been published in 19 languages, distributed worldwide,

1 listed as one of the 36 Best Books of 2011 by The Evening Standard, London, and received the
 2 Prisma Prize to the best science book of the year. Sleights of Mind appears in Books3 and Apple
 3 created unauthorized copies of it and used its to train its AI models that compete with Plaintiffs
 4 and diminish the value of their works and their labor.

5 19. Professor Macknik completed a triple-major in Psychobiology, Biology, and
 6 Psychology at the University of California, Santa Cruz in 1991. Thereafter, he completed his PhD
 7 in Neurobiology at Harvard University in 1996. He received his postdoctoral training from the
 8 Nobel Laureate Prof. David Hubel at Harvard Medical School, from 1996 to 2001. His research
 9 and writing have been featured in print in *The New York Times*, *The New Yorker*, *The Wall Street*
 10 *Journal*, *The Atlantic*, *Wired*, *The LA Chronicle*, *The Times (London)*, *The Chicago Tribune*, *The*
 11 *Boston Globe*, *Der Spiegel*, and in radio and TV shows, including *Discovery Channel's Head*
 12 *Games* and *Daily Planet* shows, *NOVA: scienceNow*, *CBS Sunday Morning*, *NPR's Science*
 13 *Friday*, and *PRI's The World*.

14 20. Professors Martinez-Conde and Macknik have multiple copyright registrations
 15 including:

Full Title	Registration No.	Date
<u>CHAMPIONS OF ILLUSION: The Science Behind Mind-Boggling Images and Mystifying Brain Puzzles</u>	TX0008595717	2/2/2018
<u>SLEIGHTS OF MIND: What the Neuroscience of Magic Reveals About Our Everyday Deceptions</u>	TX0007300820	12/27/2010

16
17
18
19
20
21
22 21. Sleights of Mind is included in the Books3 dataset. Apple copied Sleights of Mind
 23 without authorization, and used it to train its AI models. Champions of Illusion can be found in
 24 Library Genesis shadow library.

25 B. Defendant

26 22. Defendant Apple Inc. is a California corporation with its principal place of
 27 business at One Apple Park Way, Cupertino, CA 95014.
 28

C. Agents and Co-Conspirators

23. The unlawful acts alleged against Defendants in this class action complaint were authorized, ordered, or performed by the Defendants' respective officers, agents, employees, representatives, or shareholders while actively engaged in the management, direction, or control of the Defendants' businesses or affairs. The Defendants' agents operated under the explicit and apparent authority of their principals. Each Defendant, and its subsidiaries, affiliates, and agents operated as a single unified entity.

24. Various persons or firms not named as defendants may have participated as co-conspirators in the violations alleged herein and may have performed acts and made statements in furtherance thereof. Each acted as the principal, agent, or joint venture of Defendants with respect to the acts, violations, and common course of conduct alleged herein.

V. FACTUAL ALLEGATIONS

25. Apple is one of the three largest companies in the world with a \$3.8 trillion market capitalization. Apple is an electronics and media company that designs, manufactures, and sells software and hardware products, including the ubiquitous iPhone. Every second, Apple sells seven iPhones. It sells television, music, movies, podcasts, books and other media to consumers through its proprietary applications on its phones, tablets and computers. Indeed, Apple has licenses related to many of the digital book versions of the books at issue in this lawsuit, for the limited purpose of selling them to consumers through their Apple Books app. Any use of such books beyond that purpose is infringement.

26. Apple is also building commercial AI products. In January 2025 the company reported its "best quarter ever" with revenue of \$124.3 billion, twice citing Apple Intelligence in its press release announcing the same and deeming it a part of the company's "best-ever lineup of products and services." The technology is integrated across Apple's products—including iPhones—and is intended to "make[] apps and experiences even better and more personal."

27. In or around June 2024, Apple announced the development of its commercial artificial intelligence platform, called Apple Intelligence.¹ Apple Intelligence includes multiple generative-AI models and related tools and technologies. The day after Apple officially introduced Apple Intelligence the company gained more than \$200 billion in value: “the single most lucrative day in the history of the company.”

28. To train the generative-AI models that are part of Apple Intelligence, Apple first amassed an enormous library of raw training data. Apple scraped the internet with a web-crawler called “Applebot.” It also utilized training datasets that contained pirated books in the Books3 shadow library. Part of Apple’s data library includes unauthorized copies of Plaintiffs’ and Class members’ copyrighted works that were copied without authorization.

29. Apple has not attempted to pay these authors for their contributions to their commercial AI products that compete directly with them and their copyrighted works. Apple did not seek licenses to copy and use the copyrighted books it used to train to its models. Instead, it intentionally evaded payment by using books already compiled in pirated datasets.

A. How large language models work.

30. Artificial intelligence—commonly abbreviated “AI”—denotes software that is designed to algorithmically create an illusion of human reasoning or inference, often using statistical and mathematical methods.

31. The Apple Intelligence platform includes multiple AI software programs called large language models (“LLMs”) that Apple created for use in its products. Apple created these generative models for use in a wide range of AI features integrated across its platforms. An LLM is AI software designed to emit convincingly naturalistic text outputs in response to user prompts.

32. Apple’s LLM training started by amassing large quantities of raw, pre-training data. Apple admits that it sources a significant portion of the pre-training data for its models from web content crawled by Applebot, spanning hundreds of billions of links and pages, covering an extensive range of languages, locales, and topics. Apple represents that AppleBot “employs

¹ See <https://machinelearning.apple.com/research/introducing-apple-foundation-models>

1 advanced crawling strategies to prioritize high-quality and diverse content,” and that “high-
2 quality filtering plays a critical role in overall model performance.” Apple’s web crawler made
3 unauthorized reproductions of online pirated book libraries known as “Shadow Libraries”
4 including the “Books3” shadow library in which Plaintiffs’ works are included.

5 33. Apple also curates and uses training datasets that include copied shadow libraries,
6 including Books3, to train its models.

7 34. Apple then processes that pre-training data for use as model training data. Raw
8 datasets are processed and filtered in various ways in pretraining, creating derivative datasets that
9 are fed into models for training analysis. Apple applies filters to remove certain categories of
10 personally identifying information, copyright notices and license language, unwanted language,
11 profanity and unsafe material. This corpus of text is called the “training dataset.”

12 35. An LLM is “trained” by analyzing text data to extract massive amounts of
13 information about the relationships between each “token” in a textual work. A token is a small
14 string of characters, a word or sometimes a piece of a word, that serves as the base unit of AI
15 models. AI models process the relationships between each token in a work and stores that
16 information in what are called “weights.” The LLM copies and analyzes each textual work in the
17 training dataset and extracts the protected expression from it in the finest detail, on a token-by-
18 token basis. The LLM records the results of this process for each token within the model. These
19 weights are entirely and uniquely derived from the protected expression in the training dataset.
20 Generally, the more data the LLM copies during training, the better the LLM’s ability to simulate
21 the protected expression within that data as part of the LLM’s output.

22 36. Weights are the fundamental control knobs or settings for a model and determine
23 how a model evaluates new data and makes predictions. They are the core parameters for an LLM
24 and are learned during training. LLMs can have millions or even billions of weights. Weights are
25 numerical variables that set the relative importance of the features in the dataset on the output.
26 They are based entirely on the relationships between the words of copyrighted materials.

27 37. Once the LLM has copied and ingested the textual works in the training dataset
28 and converted the protected expression into stored weights, the LLM can emit convincing

1 simulations of natural written language in response to user prompts by tokenizing that prompt
 2 and using mathematics to generate the most probable string of tokens in response to that prompt.
 3 Whenever an LLM generates a response to a user prompt, it is performing a computation that
 4 relies on these stored weights recorded from copies of Plaintiffs' and Class members' works, and
 5 is imitating the protected expression ingested from the training dataset. There is no human
 6 ingenuity or creativity in any model output. It is a computer program performing mathematical
 7 operations using information it copied and extracted from its training data.

8 38. LLMs have a universal problem in “memorizing” their training data and
 9 regurgitating or reciting their training data in outputs from time to time. Apple Intelligence
 10 models are no different. During training, an LLM processes pieces of text in its training data and
 11 learns probabilistic relationships across the dataset as a whole. These learned patterns are not so
 12 high-level or abstract; rather, they are often highly specific. Training data is not so transformed by
 13 training, rather it is sometimes “memorized” by the model—encoded in some form inside the
 14 LLM's weights. Such memorized content can sometimes be “extracted” later; they can be
 15 reproduced in an LLM's outputs at generation time.²

16 39. The ability to extract verbatim portions of training data generates a copy of
 17 training data, but it also demonstrates the existence of a copy of that training data is memorized
 18 inside the model itself. The model itself also being a copy has important implications. Notably, a
 19 model—not just extracted training data—is an infringing copy of the training data it has
 20 memorized. At minimum, the copied data Apple integrates to prevent regurgitation or recitation
 21 of training data in Apple Intelligence models is an unauthorized copy of Plaintiffs' and Class
 22 members' works. Copyright law offers the destruction of infringing materials as a remedy.

23 **B. Apple trained its OpenELM models on copyrighted works.**

24 40. In April 2024, Apple first announced the availability of the OpenELM language
 25 models on its website: “[W]e release OpenELM, a state-of-the-art open language model.
 26
 27
 28

² See <https://arxiv.org/pdf/2404.12590>

1 OpenELM uses a layer-wise scaling strategy to efficiently allocate parameters within each layer of
2 the transformer model, leading to enhanced accuracy.”

3 41. The set of OpenELM language models released in April 2024 included variants
4 called OpenELM-270M, OpenELM-450M, OpenELM-1_1B, and OpenELM-3B. The main
5 difference between these models is the parameter size; a larger parameter size means the model
6 can store more tokens and weights and perform more complex tasks (requiring more computing
7 power). For instance, Apple’s OpenELM-3B language model is so named because the model
8 stores three billion (“3B”) parameters (the weights and biases) recorded from the protected
9 expression found in its training dataset.

10 42. Each OpenELM model is hosted on a website called Hugging Face, where it has a
11 “model card,” a file accompanying an AI model that typically describes the model, its intended
12 uses and limitations, its training parameters, and the training dataset used to train the model. The
13 model card for each OpenELM model states “Our pre-training dataset contains RefinedWeb,
14 deduplicated PILE, a subset of RedPajama, and a subset of Dolma v1.6, totaling approximately 1.8
15 trillion tokens.”³

16 43. Apple’s GitHub repository confirms “OpenELM was pretrained on public
17 datasets. Specifically, our pre-training dataset contains RefinedWeb, PILE, a subset of
18 RedPajama, and a subset of Dolma v1.6.”⁴

19 44. The Pile (Gao et al., 2020; Biderman et al., 2022), is a curated collection of
20 English language datasets including Books3 used that is popular for training large LLMs. It was
21 curated by a research organization called EleutherAI. Books3 is a component of The Pile.⁵
22
23
24
25

26 ³ See <https://huggingface.co/apple/OpenELM>

27 ⁴ See [https://github.com/apple/corenet/blob/main/projects/openelm/README-](https://github.com/apple/corenet/blob/main/projects/openelm/README-pretraining.md)
28 [pretraining.md](https://github.com/apple/corenet/blob/main/projects/openelm/README-pretraining.md)

⁵ See <https://arxiv.org/pdf/2201.07311v1>

45. In December 2020, EleutherAI introduced this dataset in a paper called “The Pile: An 800GB Dataset of Diverse Text for Language Modeling” (“The Pile Paper”). This paper describes the contents of Books3.⁶

Books3 is a dataset of books derived from a copy of the contents of the Bibliotik private tracker ... Bibliotik consists of a mix of fiction and nonfiction books and is almost an order of magnitude larger than our next largest book dataset ... We included Bibliotik because books are invaluable for long-range context modeling research and coherent storytelling.

46. Apple also published a paper about OpenELM (“OpenELM Paper”).⁷ In a table called “Dataset used for pre-training OpenELM,” Apple reveals that a large quantity of training

Source	Subset	Tokens
RefinedWeb		665 B
	Github	59 B
	Books	26 B
	ArXiv	28 B
	Wikipedia	24 B
	StackExchange	20 B
RedPajama	C4	175 B
PILE		207 B
Dolma	The Stack	411 B
	Reddit	89 B
	PeS2o	70 B
	Project Gutenberg	6 B
	Wikipedia + Wikibooks	4.3 B

Table 2. Dataset used for pre-training OpenELM.

data comes from the “Books” subset of a dataset called “RedPajama.”

47. Apple’s OpenELM models were trained on RedPajama-V1.⁸ RedPajama-V1 “is a publicly available, fully open, best-effort reproduction of the training data...used to train the first iteration of LLaMA family of models.” This LLaMA training dataset included Books3 section of The Pile.⁹

⁶ See <https://arxiv.org/pdf/2101.00027>

⁷ See <https://arxiv.org/pdf/2404.14619>

⁸ See <https://arxiv.org/html/2411.12372v1>.

⁹ See <https://arxiv.org/pdf/2302.13971>.

1 48. The RedPajama dataset is hosted on Hugging Face. According to the
2 documentation for the RedPajama dataset that was available there until around April 2024, its
3 “Books” component is a copy of the “Books3 dataset” that is “downloaded from Huggingface
4 [sic]” when a user runs the script that automatically assembles the RedPajama dataset. Therefore,
5 anyone who used the “Books” subset of the RedPajama dataset for training an AI model used a
6 copy of the Books3 dataset. The documentation for the RedPajama dataset does not further
7 describe the contents of Books3.

8 49. Bibliotik is one of several notorious shadow libraries, along with Library Genesis
9 (aka LibGen, Z-Library, or B-ok), Sci-Hub, and Anna’s Archive. The AI-training community has
10 long been interested in these shadow libraries because they host and distribute vast quantities of
11 unlicensed copyrighted material. For that reason, these shadow libraries violate the U.S.
12 Copyright Act.

13 50. The person who assembled the Books3 dataset, Shawn Presser, has confirmed in
14 public statements that it represents “all of Bibliotik” and contains approximately 196,640 books.

15 51. The 196,640 books in the Books3 dataset exist in .txt file format. A .txt file
16 (pronounced a “text” file) is a simple file format that stores text data without any formatting,
17 fonts, or images. Accordingly, the Books3 dataset consists of the text of the underlying 196,640
18 books.

19 52. By using the entire text, Apple made and used copies of each book in Books3 to
20 store as reference data and to train its OpenELM models and develop other analytic and filtering
21 tools.

22 53. Plaintiffs’ copyrighted works are among the works in the Books3 dataset.

23 54. Until October 2023, the Books3 dataset was available from Hugging Face. At that
24 time, the Books3 dataset was removed with a message that it “is defunct and no longer accessible
25 due to reported copyright infringement.”

26 55. Presser himself has acknowledged that “we almost didn’t release the data sets at
27 all because of copyright concerns.”
28

56. Before October 2023, anyone who used the “Books” subset of the RedPajama dataset for training necessarily made a copy of the Books3 dataset. Based on Apples’s own information revealed in the OpenELM research paper and on its model card on Hugging Face, this includes Apple.

57. Apple has admitted training its OpenELM large language models on a copy of the “Books” subset of the RedPajama dataset, as well as a deduplicated version of The Pile, each of which in turn contains a copy of the Books3 dataset. Therefore, Apple trained its OpenELM models on a copy of Books3, a known body of pirated books.

58. Because Plaintiffs’ copyrighted book is part of Books3, Apple copied in its entirety without authorization, and trained OpenELM, on one or more copies of the Plaintiffs copyrighted works and directly infringed Plaintiffs’ copyrights along with the copyrights of the Class.

C. Apple trained its Foundation Language Models on copyrighted works.

59. Apple announced their two Apple Intelligence Foundation Language Models in June 2024. These models were described in a paper of the same name released by Apple on July 29, 2024 (the “FLM Paper”).¹⁰

60. The adjective *foundation* is commonly used to describe AI models that have broad capabilities to perform a wide variety of tasks. Consistent with this, Apple describes its Foundation Language Models as “highly capable in tasks like language understanding, instruction following, reasoning, writing, and tool use ... These foundation models are at the heart of Apple Intelligence.”

61. The FLM Paper emphasizes the special importance of a foundation model’s capacity to write: “Writing is one of the most critical abilities for large language models to have, as it empowers various downstream use cases such as changing-of-tone, rewriting, and summarization.”

62. In the FLM Paper, Apple identifies two separate foundation language models: AFM-server and AFM-on-device. The AFM-server model is a larger model that is intended for

¹⁰ See <https://arxiv.org/pdf/2407.21075>

1 use through an Apple-operated cloud service called Private Cloud Compute. The AFM-on-
2 device model, by contrast, is intended to be small enough to be used directly on Apple devices
3 (e.g., iPhones and laptops). According to the FLM Paper, the AFM-on-device model is
4 “initialize[d] ... from a pruned 6.4B model (trained from scratch using the same recipe as AFM-
5 server.)”. This means that both the AFM-server and AFM-on-device models are trained on the
6 same corpus of training data.

7 63. In the FLM Paper, Apple reveals three sources of training data: “data we have
8 licensed from publishers, curated publicly available or open-sourced datasets, and publicly
9 available information crawled by our web-crawler, Applebot.”

10 64. In describing the first source of training data—“data we have licensed”—Apple
11 says only that it “identif[ied] and license[d] a **limited amount** of high-quality data from
12 publishers” (emphasis added). In addition to being comparatively “limited” in quantity, Apple
13 does not use this licensed data during the main phase of training the Foundation Language
14 Models—what Apple calls “core pre-training”—but during a subsequent phase called “continued
15 pre-training.”

16 65. As to its second source of training data—“publicly-available or open-sourced
17 datasets”—Apple does not elaborate on the specific datasets used, saying only, “We evaluated
18 and selected a number of high-quality publicly-available datasets with licenses that permit use for
19 training language models.” Apple then we filtered the datasets to remove personally identifiable
20 information, copyright management information, and other undesired text before including them
21 in the pre-training mixture.

22 66. In the parlance of AI training datasets, Apple’s phrase “publicly available” is one
23 commonly used to falsely conjure up the idea of works made publicly available by the author. In
24 practice, “publicly available” means works that can be downloaded somewhere from the public
25 internet, which contains a vast number of copyrighted works by authors who have not granted a
26 license for reproduction. There is a name for this kind of copying: copyright infringement. There
27 is also a name for the infringing copies: pirated works.

1 67. For instance, Meta Platforms also trained its Llama language models on Books3,
2 which it described as a “publicly available dataset for training large language models” despite the
3 fact that none of the authors whose works appear in Books3 ever consented to having their works
4 included. Books3 was “publicly available” only in the limited sense that at one time, it could be
5 acquired by anyone with an internet connection.

6 68. Similarly, in the context of AI training datasets, Apple’s phrase “open source” is
7 commonly used to falsely conjure up the idea of works made available by the author under a
8 permissive copyright license (e.g., a Creative Commons license). In practice, what it really means
9 is someone other than the author made curated copies of copyrighted works freely available,
10 without the author’s permission. These are just pirated works included in a curated dataset that
11 the pirate claims is open-source.

12 69. For instance, EleutherAI, the group that created The Pile—the dataset that
13 included Books3—described it as a “diverse, open source language modelling data set” even
14 though the copyrighted works in the Books3 portion were included without authors’ consent.
15 Only the copyright owner can offer their copyrighted work to the public under an open-source
16 license. A third party cannot usurp that right.

17 70. Therefore, when Apple says that a major source of the training data for its
18 Foundation Language Models is “publicly-available or open-sourced datasets,” they are talking
19 about curated datasets like RedPajama. Because Books3 has been described by people in the AI
20 industry as a “publicly available” or “open-sourced” dataset, and because Apple already had a
21 copy of Books3 that it had used for training its OpenELM models, Apple’s reference to “publicly-
22 available or open-sourced datasets” likely includes Books3, and that Apple therefore included
23 Books3 in the training dataset for its Foundation Language Models.

24 71. Plaintiffs’ and Class members’ copyrighted works are part of Books3. It follows
25 that Apple trained its Foundation Language Models on one or more copies of each book therein,
26 thereby directly infringing the copyrights of the Plaintiffs and the Class. Apple has created a
27 permanent AI training data library containing copies of all these “publicly-available or open-
28 sourced datasets” in expectation of training future models.

1 72. As to its third source of training data—web pages crawled by Applebot—Apple
2 says, “we crawl publicly available information using our web crawler, Applebot ... and respect the
3 rights of web publishers to opt out of Applebot.” Applebot has been crawling the web since
4 approximately mid-2015. Around June 2024, Apple revealed that it was using Applebot-scraped
5 data for training its AI models. In response to this disclosure, by August 2024, numerous major
6 commercial web publishers had chosen to opt out of Applebot training.

7 73. But Apple’s Foundation Language Models had necessarily been trained well
8 before the release of the FLM Paper describing them in July 2024. For that reason, Apple’s
9 disclosure in June 2024 that it was using Applebot data to train language models came too late for
10 any of these opt-outs to matter. Apple had already scraped the data and trained language models
11 with it. On information and belief, Apple has retained copies of all AppleBot data scraped before
12 this wave of opt-outs, in expectation of training future models, as part of its AI training-data
13 library.

14 74. In the FLM Paper, Apple says that Applebot pages are “processed by a pipeline
15 which performs quality filtering ... using heuristics and model-based classifiers.” In this context,
16 the term “model-based classifier” refers to a separate AI model that has been trained to
17 algorithmically rate the quality of scraped web pages. These model-based classifiers are
18 themselves trained on datasets that include unlicensed copyrighted works.

19 75. In a November 2024 paper by George Wukoson and Joey Fortuna called “The
20 Predominant Use of High-Authority Commercial Web Publisher Content to Train Leading
21 LLMs,”¹¹ the authors studied LLM training datasets made by algorithmically filtering scraped
22 web pages. The authors concluded that such “datasets are disproportionately composed of high-
23 quality content owned by commercial publishers of news and media websites.” In turn, this
24 material is often covered by registered copyrights. Thus, the part of Apple’s training dataset that
25 comes from filtered Applebot pages includes copyrighted works from commercial news and media
26 websites.

27
28

¹¹ See <https://papers.ssrn.com/sol3/Delivery.cfm/5009668.pdf>

1 76. The shadow libraries that host millions of unlicensed copyrighted books are also
2 part of the “publicly available information” reachable by a web scraper like Applebot. Hence, on
3 information and belief, part of Apple’s training-data library is sourced from shadow libraries via
4 Applebot directly.

5 77. Apple obscures the training datasets for its Apple Intelligence Foundational
6 Language Models to blur its use of copyrighted materials. Apple’s decision not to disclose the
7 training datasets for Apple Intelligence stems in part from the fact that Apple was the subject of
8 negative press for using a subset of data from The Pile containing captions from thousands of
9 YouTube videos.

10 78. The “curated publicly available or open-sourced datasets” that Apple copied for
11 the training datasets for Apple Intelligence contain copyrighted material, including Plaintiffs’
12 copyrighted works. Such use of datasets with copyrighted works would be consistent with Apple’s
13 process for training its OpenELM model. Apple described its training data for OpenELM,
14 including data from The Pile, as “public datasets.” But the “public” nature of a dataset does not
15 mean that the data collected in the dataset was obtained lawfully or that the party providing
16 copies of the dataset has authority to extend a valid license to use the underlying copyrighted
17 works.

18 79. There are numerous examples of publicly reported AI licensing deals. Myriad
19 licensing systems have been launched and are continuing to develop, including the Copyright
20 Clearance Center’s collective AI licensing scheme and the Created by Humans licensing
21 platform. Further, several AI data set licensing companies have formed a trade group called the
22 Dataset Providers Alliance. Currently, some researchers estimate, the AI training license market
23 is valued at approximately \$2.5 billion; within a decade, it may close in on nearly \$30 billion.

24 80. Apple itself understands the value of copyrighted works and the market that exists
25 for paying creators to use their works for training. For instance, it struck an agreement with
26 Shutterstock to “use hundreds of millions of images, videos and music files” valued between an
27 estimated \$25 to \$50 million.

1 81. Similarly, Apple has contacted news organizations like Condé Nast, NBC News,
 2 and IAC to license news article archives. Nonetheless, Apple has not compensated Plaintiffs' and
 3 Class members whose works it copied and used in trainings its models.

4 82. Furthermore, Apple is reportedly exploring a paid tier for users of its Apple
 5 Intelligence products. Doing so might be in effort to offset the costs of its steep investments in
 6 building Apple Intelligence. Analysts contend that Apple Intelligence could add \$4 trillion to the
 7 company's market capitalization.

8 **D. Apple's conduct impairs the market for Plaintiffs' and Class members' works.**

9 83. Apple has neither paid nor sought permission from Plaintiffs for the use of their
 10 copyrighted works. Instead, Apple downloaded, scraped, or otherwise copied vast quantities of
 11 copyrighted works—including illegally compiled datasets such as Books3—that included
 12 Plaintiffs' works. Apple has diminished the value of Plaintiffs' intellectual property by making
 13 unauthorized copies and derivative works for use in training their Apple Intelligence models.
 14 Apple further deprived Plaintiffs of the revenue that would have been generated had Apple
 15 approached Plaintiff or their licensing agents directly to license copies of their works for use in AI
 16 training. Furthermore, compiling private libraries sourced from illegally compiled datasets for AI
 17 training purposes may “lead to a loss of sales” by “harm[ing] the market for access to those
 18 works.”

19 84. Apple's unauthorized use of Plaintiffs' copyrighted works creates a risk of market
 20 dilution of the works themselves and of Plaintiff's professional writing. Works generated using
 21 Apple Intelligence will inevitably start competing with Plaintiffs' copyrighted and other works
 22 and ultimately dilute royalty pools and professional writing opportunities as AI-generated output
 23 increasingly floods the market. Already, “low-quality sham ‘books’” have begun overwhelming
 24 the market as scammers generate “unauthorized ‘biographies’ of authors that are simply AI-
 25 generated rehashings of their lives, often based on autobiographical works.” Other scams include
 26 “companion books” that summarize the key points from the original novel, with “little to no
 27 original analysis or commentary and are meant only to confuse consumers and skim sales off of
 28 the real books.” These works have already entered book marketplaces like Amazon.

85. Apple’s unauthorized use of Plaintiffs’ copyrighted works to train Apple Intelligence models has caused and threatens to cause substantial harm to the actual potential markets for those works. Plaintiffs and similarly situated creators previously licensed their work for their own commercial uses. Apple’s conduct has disrupted this traditional market and impaired the emergence of lawful licensing regimes by obtaining and exploiting authors’ works without consent or compensation.

86. Apple’s models generate outputs by automated computer operation that substitute for the kinds of expressive written work that Plaintiffs are hired to produce, potentially diminishing demand for books and human-produced stories. Plaintiffs face potential ongoing harms through lost publication opportunities and reduced recognition, and sales among other harms.

VI. CLASS ACTION ALLEGATIONS

87. The “Class Period” as defined in this Complaint begins at least three years before the date of this complaint’s filing and runs through the present.

88. Apple engaged in a course of conduct common to all class members in infringing Plaintiffs’ and Class members’ works including:

- a. Apple reproduced all class member books in the Books3 dataset without authorization to train the OpenELM models;
- b. Apple further reproduced all class member books in the Books3 dataset without authorization to train its Foundation Language Models;
- c. Apple made unauthorized copies and prepared of unauthorized derivative works of all class member books in the course of pre-training for those models;
- d. Apple made and used unauthorized copies of datasets that included class members’ unlicensed copyrighted works to train classifier modes; and
- e. Apple retained copies of all the training data it has gathered and processed thus far without authorized, in the form of a private data library for potential use in future models—an AI training data library —that includes the Books3 dataset, which, in turn, includes the Plaintiffs’ and Class members’ infringed works.

1 **89. Class definition.** Plaintiffs bring this action for damages and injunctive relief as a
2 class action under Federal Rules of Civil Procedure 23(a), 23(b)(2), and 23(b)(3), on behalf of the
3 following Class:

4 All legal or beneficial owners of a registered copyright for any work Apple copied
5 without authorization or ingested for training or otherwise developing its
6 OpenELM Models and Foundation Models during the relevant time period, which
7 was registered with the United States Copyright Office within five years of the
8 work's publication and which was registered with the United States Copyright
9 Office before being trained on by Apple, or within three months of publication.
Excluded from the Classes are the Defendant, its subsidiaries and affiliates, officers,
executives, and employees; Defendant's attorneys in this case, federal government
entities and instrumentalities, states or their subdivisions, and all judges and jurors
assigned to this case.

10 90. Plaintiffs reserve the right to modify or amend the definition of the proposed Class
11 before the Court determines whether certification is appropriate.

12 91. The Class members are so numerous and geographically dispersed that individual
13 joinder of all Class members is impracticable. Moreover, given the costs of complex antitrust
14 litigation, it would be uneconomical for many Plaintiffs to join their individual claims. The exact
15 number of Class members is currently unknown to Plaintiffs, as this information is in Defendant's
16 exclusive control. On information and belief, there are more than several thousand members in
17 the Class across the United States. Accordingly, joinder of all Class members in prosecuting this
18 action is impracticable.

19 92. The Class can be identified, in part, through tools that allow a user to search for
20 web domains included in the RedPajama dataset, The Pile dataset, or other datasets used to train
21 one or more of the Apple Intelligence models.

22 93. The Class can further be identified by analyzing the training data that Apple used
23 for both its OpenELM Models and Apple Foundation Models.

24 94. ISBN numbers can be used to identify copyrighted books in Apple's training data.

25 95. Plaintiffs' claims are typical of the claims of Class Members because Plaintiffs and
26 all members of the Class were damaged by the same course of conduct of Defendant. Further, the
27 relief sought is common to all Class members.
28

1 96. Plaintiffs will fairly and adequately protect and represent the interests of the Class.
2 The interests of the Plaintiffs are aligned with, and not antagonistic to, those of the other Class
3 members. Further, Plaintiffs have retained competent counsel who are experienced in litigating
4 federal class actions and other complex class action litigation involving copyright infringement in
5 the context of training AI models.

6 97. Numerous questions of law and fact are common to each Class member arising
7 from Defendant's conduct, including:

- 8 a. Whether Plaintiffs' and Class members' works were included in the training datasets
9 used by Apple to train its Apple Intelligence product, including the RedPajama and
10 Pile datasets Defendant used;
- 11 b. Whether Apple's inclusion of Plaintiffs' and Class members' works in their training
12 datasets constituted or required the works' unauthorized reproduction by Apple;
- 13 c. Whether Apple lacked authorization to reproduce or make copies of Plaintiff's and
14 Class members' works;
- 15 d. Whether Apple violated Plaintiff's and the Class members' copyrights when it
16 downloaded copies of Plaintiff's copyrighted works and used them in training its
17 OpenELM and Foundation Models;
- 18 e. Whether Apple violated Plaintiff's' and Class members' copyrights when it
19 downloaded copies of Plaintiff's copyrighted works without authorization and used
20 them in its Apple Intelligence products;
- 21 f. Whether Apple violated Plaintiff's and Class members' copyrights by creating
22 unauthorized copies or derivative works based on their copyrighted works in pre-
23 training and training of the OpenELM and Foundation Models;
- 24 g. Whether Apple violated Plaintiff's and Class members' copyrights by creating
25 unauthorized copies or derivative works based on their copyrighted works in creating
26 the OpenELM and Foundation Models themselves, or the filters they employ to
27 prevent the regurgitation of copyrighted works in model outputs;
- 28 h. Whether Apple violated Plaintiff's and Class members' copyrights by creating

1 unauthorized copies or derivative works based on their copyrighted works in pre-
 2 training and training of the OpenELM and Foundation Models;

3 i. Whether this Court should enjoin Defendant from engaging in the unlawful conduct
 4 alleged herein or provide other equitable relief;

5 j. Whether any affirmative defense excuses Defendant's conduct, including the fair use
 6 doctrine; and

7 k. Whether Apple's infringement was willful; and

8 l. the appropriate measure of damages.

9 98. These and other questions of law and fact are common to the Class and
 10 predominate over questions affecting Class members on an individual basis. Damages pose no
 11 obstacle for class certification as statutory damages are available to all Class members pursuant to
 12 17 U.S.C. § 504.

13 99. Defendant has acted on grounds generally applicable to the Class. A class action is
 14 superior to alternatives for the fair and efficient resolution of this controversy. Allowing the
 15 claims to proceed on a class basis will eliminate the possibility of repetitive litigation. Further,
 16 injunctive relief is appropriate with respect to the entire Class. The alternative of separate actions
 17 by individual Class members risks inconsistent adjudications and is an inefficient use of limited
 18 judicial resources.

19 100. Plaintiffs know of no difficulty to be encountered in the maintenance of this action
 20 that would preclude its maintenance as a class action.

21 **VII. CLAIMS**

22 **COUNT ONE** 23 **DIRECT COPYRIGHT INFRINGEMENT** 24 **(17 U.S.C. § 501)**

25 101. Plaintiffs incorporate by reference all other allegations in this complaint.

26 102. As the owners of the registered copyrights in their copyrighted books and other
 27 copyrighted works, Plaintiffs and Class members hold the exclusive copyrights in those works
 28 under 17 U.S.C. § 106, including the exclusive right to: (1) reproduce the copyrighted work in

copies; (2) prepare derivative works based upon the copyrighted work; (3) distribute copies of the copyrighted work; (4) perform the copyrighted work; and (5) display the copyrighted work to the public. Plaintiffs and the Class members never authorized Apple to make copies of their copyrighted books and other copyrighted works, prepare derivative works, publicly display copies or derivative works, or distribute copies or derivative works, or exploit any other right exclusively reserved to Plaintiffs and the Class members under the U.S. Copyright Act.

103. Apple made all its copies of Plaintiffs' copyrighted books and other copyrighted works without Plaintiff's or Class members' permission, violating their exclusive rights under the U.S. Copyright Act. Indeed, "the person who copies the textbook from a pirate website has infringed already, full stop." *Bartz et al. v. Anthropic*, No. C 24-05417 WHA, 2025 WL 1741691 at *11 (N.D. Cal. June 23, 2025). Regardless of how Apple uses the works in its private training-data library in the future, this cannot negate the initial copying of works sourced from shadow libraries infringed on Plaintiff's and Class members' exclusive rights.

104. Plaintiffs and Class members have been injured by the Apple's direct copyright infringement of the their copyrighted works. Plaintiffs and Class members are entitled to statutory damages, actual damages, restitution of profits, and other remedies provided by law or in equity.

VIII. PRAYER FOR RELIEF

Plaintiffs demand judgment on their behalf and on behalf of the Class against each Defendant as follows:

- a. Allowing this action to proceed as a class action, with Plaintiffs serving as Class Representatives, and with Plaintiff's counsel as Class Counsel;
- b. Awarding Plaintiffs and the Class statutory damages, compensatory damages, restitution, disgorgement, and any other relief that may be permitted by law or equity;
- c. Permanently enjoining Defendant from the unlawful, unfair, and infringing conduct alleged herein;
- d. Ordering destruction under 17 U.S.C. § 503(b) of all Apple Intelligence or other LLM models and training sets that incorporate Plaintiff's and Class members' works;

- 1 e. An award of costs, expenses, and attorneys' fees as permitted by law; and
2 f. Such other or further relief as the Court may deem appropriate, just, and equitable.

3 **IX. DEMAND FOR JURY TRIAL**

4 Plaintiffs demand a jury trial for all claims.

5
6 Dated: October 9, 2025

Respectfully Submitted,
JOSEPH SAVERI LAW FIRM, LLP.

7
8 By: /s/ Joseph R. Saveri
9 Joseph R. Saveri

10 Joseph R. Saveri (State Bar No. 130064)
11 JOSEPH SAVERI LAW FIRM, LLP
12 601 California Street, Suite 1505
13 San Francisco, California 94108
14 Telephone: (415) 500-6800
15 Facsimile: (415) 395-9940
16 Email: jsaveri@saverilawfirm.com

17 *Attorneys for Plaintiffs*
18
19
20
21
22
23
24
25
26
27
28