

SUSMAN GODFREY L.L.P.

A REGISTERED LIMITED LIABILITY PARTNERSHIP
1000 LOUISIANA
SUITE 5100
HOUSTON, TX 77002-5096
(713) 651-9366
FAX (713) 654-6666
WWW.SUSMANGODFREY.COM

401 UNION STREET
SUITE 3000
SEATTLE, WA 98101-2668
(206) 516-3880

ONE MANHATTAN WEST
NEW YORK, NY 10001-8602
(212) 336-8330

1900 AVENUE OF THE STARS
SUITE 1400
LOS ANGELES, CA 90067
(310) 789-3100

JUSTIN A. NELSON
DIRECT DIAL 713-653-7895

E-MAIL JNELSON@SUSMANGODFREY.COM

November 18, 2025

VIA ECF

Hon. Ona T. Wang
Daniel Patrick Moynihan
United States Courthouse
500 Pearl St.
New York, NY 10007

RE: *In re OpenAI, Inc. Copyright Infringement Litigation* (No. 1:25-md-03143) This document relates to the following Class Cases: Case No. 1:23-cv-08292, Case No. 1:23-cv-10211, Case No. 1:24-cv-00084, Case No. 1:25-cv-03291, Case No. 1:25-cv-03482, Case No. 1:25-cv-03483.

Dear Judge Wang,

Plaintiffs seek to inspect a limited set of OpenAI models—known as “intermediate” and “base” models—as part of their work assessing the technical nature of the training process, the value of certain training data, and the internal workings of OpenAI’s LLMs.

These assessments are relevant to fair use factors one (transformativeness) and four (market harm); the models in question are specific versions of the GPT-models which Judge Stein has already held are in-scope. There is no undue burden associated with the inspections Plaintiffs seek, as OpenAI has previously made them available to both internal and external researchers and, in some cases, paying customers. Yet OpenAI has categorically refused to permit Plaintiffs to inspect these models. The Court should order that OpenAI make these models available for inspection.

I. Base Models:

As used herein, a “base” model refers to an LLM after it has undergone the pre-training phase of the LLM training process, but before it has undergone the various modifications and filters—generally known as “mid-training,” “post-training,” or “fine-tuning”—which OpenAI may implement later. Base models generally are not governed by these later modifications and so

November 18, 2025

Page 2

may be even more capable of large-scale regurgitation of material memorized during the pre-training process. Plaintiffs seek to inspect the base model versions of all models which Judge Stein held were in-scope, e.g., for OpenAI's GPT-4 model, Plaintiffs seek to inspect the model OpenAI internally and externally labels as "GPT-4 Base."

OpenAI routinely makes base models available to outside researchers through its "Researcher Access Program." *See, e.g.*, Tweet from Sam Altman, <https://x.com/sama/status/1638576434485825536>



See also Is In-Context Learning Sufficient for Instruction Following in LLMs, Zhao et al, p. 11, <https://arxiv.org/pdf/2405.19874v3> ("We are grateful to OpenAI for providing us API credits and access to GPT-4-Base as part of the OpenAI Researcher Access Program."); *Me, Myself, and AI: The Situational Awareness Dataset (SAD) for LLMs*, Laine et al, pp. 3, 7, https://proceedings.neurips.cc/paper_files/paper/2024/file/7537726385a4a6f94321e3adf8bd827e-Paper-Datasets_and_Benchmarks_Track.pdf (noting use of GPT-4 base and thanking OpenAI for access). In the case of GPT-3.5 (the model that powered ChatGPT at launch), OpenAI previously offered the "base" version of the model (also known as "code-davinci-002") to customers via its API. *See, e.g.*, <http://web.archive.org/screenshot/https://platform.openai.com/docs/model-index-for-researchers> (archive of <https://platform.openai.com/docs/model-index-for-researchers> noting that code-davinci-002 is a GPT-3.5 base model); <https://archive.ph/XDCN0> (archive of

November 18, 2025

Page 3

<https://beta.openai.com/docs/model-index-for-researchers> noting code-davinci-002 is a GPT-3.5 base model).

A. Base Model Behavior is Highly Relevant to the Parties' Dispute Over the Transformativeness of the Training Process Under Fair Use Factor One

The pre-trained “base” model of, e.g. GPT-4, is the raw, unfiltered model, on which later model performance rests (hence the moniker “base” models). If, as Class Plaintiffs explained at the tech tutorial, the training process is accurately characterized as “compression,” the pre-trained “base” model will be more likely to regurgitate than its post-trained counterpart. Third-party research has suggested as much. *See, e.g., CopyBench: Measuring Literal and Non-Literal Reproduction of Copyright-Protected Text in Language Model Generation*, Chen et al, p. 7 Table 2, <https://arxiv.org/pdf/2407.07087> (comparing “copying and utility of pre-trained base LMs” and showing higher literal copying for large base models).

Testing these base models for regurgitation, memorization, and copyright-infringing behavior is thus relevant to the parties’ dispute over how transformative the training process is and what relationship the final models bear to their training data. In other words: both of these questions are relevant to the inquiry regarding transformativeness under fair use factor one. The fact that OpenAI must go to extreme lengths to modify this default behavior in its base models demonstrates that the pre-training process itself is not transformative. Documenting that default behavior is what Plaintiffs inspection seeks to do.

In the parties’ conferral on this issue, OpenAI claimed that base models are not relevant because intermediate copies are not infringing. First, that is an incorrect statement of the law generally, *see, e.g., Am. Geophysical Union v. Texaco Inc.*, 60 F.3d 913, 924 (2d Cir. 1994) (rejecting argument that making photocopies in the course of performing scientific research was transformative), and the law regarding intermediate copying of text in generative AI products in particular, *see, e.g., Thomson Reuters Enter. Ctr. GmbH v. Ross Intel. Inc.*, 765 F. Supp. 3d 382, 398–99 (D. Del. 2025), *motion to certify appeal granted*, No. 1:20-CV-613-SB, 2025 WL 1488015 (D. Del. May 23, 2025) (holding intermediate copying cases did not apply where the challenged conduct did not involve intermediate copies of computer code). Second, OpenAI’s legal contention that it does not infringe obviously cannot control; the whole point of the discovery process (and the subsequent summary judgment proceedings) is to test that contention. Third, OpenAI has previously made such base models available to third-parties, including, in at least the cases of GPT-3 and GPT-3.5, paying customers. These are thus versions of in-scope models and their behavior is just as relevant as the other models which Plaintiffs will test.

B. OpenAI Cannot Show Any Burden Associated With Making Such Models Available for Inspection

November 18, 2025
Page 4

There is no significant burden associated with permitting Class Plaintiffs to inspect these base models: as noted above OpenAI routinely makes them available to outside researchers through its “Researcher Access Program.” And in the case of GPT-3.5 (the model that powered ChatGPT at launch), and GPT-3 OpenAI previously offered the “base” models to customers via its API.

To the extent OpenAI claims any burden from the fact that these models have been “deprecated,” that argument is foreclosed by the scope of the models in suit, which already include a large number of “deprecated” models, e.g. GPT-3, and 3.5, and OpenAI has previously made such models available to inspect.

II. Intermediate Models:

The situation is similar with respect to so-called “intermediate models” (sometimes referred to as “intermediate model checkpoints”). While the term “base” models refers to the models specifically after they have finished the pre-training process, the phrase “intermediate models” refers more generally to models at any stage of their training process prior to deployment. These intermediate models include, e.g., the models after they have been pretrained on certain datasets (but not others), or models which have undergone certain post-training or fine-tuning steps but not others. In addition to its specific refusals regarding base models, OpenAI has categorically refused to make any other intermediate models available for inspection for the same reasons discussed above.

A. Intermediate Model Performance is Highly Relevant to Fair Use Factors 1 and 4

Intermediate models are relevant for precisely the same reasons as base models, namely the performance of an intermediate model which has been, for example, trained on Plaintiffs works vs. a model which has not (or not yet) trained on Plaintiffs works is relevant to assessing the nature of the training process under fair use factor one. Their performance is also relevant to assessing the commercial and technical viability of models trained on certain data (but not Plaintiffs’ copyrighted books) and the value of Plaintiffs works to the models’ ultimate performance in general under fair use factor four.

B. There is No Significant Burden Associated With Plaintiffs’ Inspections of Intermediate Models

While OpenAI’s researchers have access to intermediate models for testing and development, OpenAI also makes intermediate models available externally. OpenAI makes clear that, as part of the model training process, it exposes internal intermediate model checkpoints to external researchers for so-called “red-teaming.” *See, e.g.,* <https://control-plane.io/case-studies/openai-red-teaming/> (“[t]esting began with live access to model checkpoints and progressively evolved as safeguards were adapted to mitigate discovered risks.”);

November 18, 2025

Page 5

<https://openai.com/index/gpt-4o-system-card/> (“Deployment preparation was carried out via exploratory discovery of additional novel risks through expert red teaming, starting with early checkpoints of the model while in development . . . Red teamers had access to various snapshots of the model at different stages of training . . .”). This external access is often done through APIs, the very same format which Plaintiffs will be using to perform all model inspections. See <https://cdn.openai.com/papers/openais-approach-to-external-red-teaming.pdf> (noting the variety of interfaces through which such access is provided, including “direct API access”). There is thus no categorical burden associated with making such model available to Plaintiffs and their experts.

If OpenAI has concerns in a given case, Plaintiffs are willing to work to address those concerns (and are confident they can be addressed on a case-by-case basis). However, particularly given OpenAI’s own public practices of exposing intermediate models to external researchers, there can be no categorical bar to Plaintiffs’ inspection of these models.

Sincerely,

/s/Justin A. Nelson

Susman Godfrey, LLP
Interim Lead Class Counsel