

1 MAYER BROWN LLP
2 EDWARD D. JOHNSON
3 *WJohnson@mayerbrown.com*
4 ANTHONY J. WEIBELL
5 *AWeibell@mayerbrown.com*
6 KRISTIN W. SILVERMAN
7 *KSilverman@mayerbrown.com*
8 ELSPETH V. HANSEN
9 *EHansen@mayerbrown.com*
10 3000 El Camino Real, Suite 300
11 Palo Alto, CA 94306
12 Tel.: (650) 331-2000
13 Fax: (650) 331-2060

9 ANKUR MANDHANIA (SBN 302373)
10 *AMandhania@mayerbrown.com*
11 575 Market Street, Suite 2500
12 San Francisco, CA 94105
13 Tel.: (415) 874-4230
14 Fax: (415) 331-2060

13 [additional counsel listed on signature page]

14 Attorneys for Defendants
15 OPENAI FOUNDATION, OPENAI OPCO, LLC,
16 OPENAI HOLDINGS, LLC, and SAMUEL ALTMAN

16 **SUPERIOR COURT OF THE STATE OF CALIFORNIA**

17 **COUNTY OF SAN FRANCISCO**

18 MATTHEW RAINE, et al.,

19 Plaintiffs,

20 vs.

21 OPENAI, INC., et al.,

22 Defendants.

Case No. CGC-25-628528

**DEFENDANTS' ANSWER TO
AMENDED COMPLAINT AND
DEMAND FOR JURY TRIAL**

**PUBLIC VERSION OF DOCUMENT
LODGED PROVISIONALLY UNDER
SEAL**

First Amended Complaint Filed:
October 22, 2025

- 1
- 2
- 3
- 4
- 5
- 6
- 7
- 8
- 9
- 10
- 11
- 12
- 13
- 14
- 15
- 16
- 17
- 18
- 19
- 20
- 21
- 22
- 23
- 24
- 25
- 26
- 27
- 28

2
3
4
5
6
7
8

9

10
11

12

- 13
14
15
16
17
18
19
20
21
22
23
24
25
26
27

28

1 the same time, it is essential that the important issues raised by this litigation be addressed
2 accurately and grounded in facts. Defendants are committed to identifying and understanding the
3 facts, in order to continue to build safe and beneficial technology, and to further OpenAI's mission
4 of ensuring that artificial general intelligence benefits all of humanity.

5 **OPENAI AND SAFETY**

6 6. Artificial intelligence (AI) is a critically important and transformative technology
7 that can have—and already has had—significant positive impact on society.

8 7. OpenAI operates the ChatGPT artificial intelligence service (“ChatGPT”), which is
9 used by hundreds of millions of people worldwide to improve their daily lives, while also helping
10 to advance scientific discovery and medical research.

11 8. OpenAI is committed to the safety of its users, and is the leader in developing safe
12 and beneficial artificial intelligence. OpenAI has pioneered many of the foundational safety
13 practices used across the industry. This is core to OpenAI's mission to ensure that artificial general
14 intelligence benefits all of humanity.

15 9. Safety is a focus of every team at OpenAI. OpenAI employs leading safety experts
16 and researchers who are dedicated to all aspects of safety, ranging from shaping model behavior to
17 red-teaming and feedback to iterative deployment.

18 10. Safety is built into every step of OpenAI's model development process. That process
19 starts with model design and pre-training to teach the model language and intelligence and
20 continues through post-training to align the model to safety principles so that it provides helpful
21 and safe answers. OpenAI applies safety principles at every stage of model training, implementing
22 system-level guardrails and investing in long-term safety research. OpenAI also conducts extensive
23 pre-deployment evaluations, where the model goes through safety evaluations, including red-
24 teaming—all of which occur before the model is released to the public. OpenAI collaborates with
25 external experts, trusted partners, and researchers to stress test models and gather feedback, and has
26 developed a Preparedness Framework to test and evaluate catastrophic risk areas. Every model
27 developed by OpenAI, including the model underlying ChatGPT when it was used by Adam Raine,
28 undergoes that state-of-the-art safety process (including with respect to suicide and self-harm) prior

1 to deployment.

2 11. For each model, OpenAI voluntarily publishes information related to safety
3 evaluations in the model's system card and on its Safety Evaluations Hub. The system card for the
4 model used by Adam Raine, GPT-4o, details OpenAI's thorough safety evaluation of the model for
5 potential risks and its implementation of safeguards based on these tests, including its work with
6 more than 100 external red teamers to test the model. The system card also demonstrates that the
7 model passed OpenAI's rigorous and state-of-the-art tests.

8 12. Once a model is released, OpenAI continues its safety work through continuous
9 post-deployment monitoring and improvement. Through its iterative deployment approach,
10 OpenAI gradually releases its models in controlled conditions, using AI tools and people to
11 continually monitor for misuse or violations of usage policies, and undergoes extensive safety
12 committee reviews. For example, OpenAI updated GPT-4o on an ongoing basis from the time of
13 its release throughout the time that Adam Raine used it, and continues to update its models. In
14 addition, OpenAI continually researches and deploys new approaches to strengthen safeguards in
15 response to real-world usage. OpenAI also publishes its safety research to set the industry standard
16 and improve safety practices across the industry.

17 **CHATGPT AND SAFETY**

18 13. OpenAI has developed a stack of layered safeguards into ChatGPT. For example,
19 since early 2023, OpenAI's models have been trained to not provide self-harm instructions and to
20 shift into supportive, empathic language. For example, if someone writes that they want instructions
21 on how to hurt themselves, ChatGPT is trained to not comply and instead acknowledge their
22 feelings and steer them toward help. Additionally, responses that OpenAI's classifiers identify as
23 contrary to the models' safety training are automatically blocked, with stronger protections for
24 minors and logged-out users. During very long sessions, ChatGPT nudges people to take a break.

25 14. ChatGPT is also designed to refer people whose prompts express suicidal intent to
26 real-world resources like professional help. For example, in the United States, ChatGPT refers
27 people to 988 (the national suicide and crisis hotline). This logic is built into model behavior.
28 OpenAI is working closely with 90+ physicians across 30+ countries—psychiatrists, pediatricians,

1 and general practitioners—and has convened an advisory group of experts in mental health, youth
2 development, and human-computer interaction to ensure its approach reflects the latest research
3 and best practices.

4 15. OpenAI also provides users with information regarding how its products are
5 designed to behave, how users must interact with its products, and how the products should respond
6 to users that are under 18. These disclosures educate the public (including OpenAI users) on how
7 to use OpenAI products safely and properly.

8 16. OpenAI makes several such disclosures to its users. First, OpenAI has published
9 several articles explaining how ChatGPT works, including a FAQ page entitled “Does ChatGPT
10 tell the truth” which explains why users should not rely on LLM output—including ChatGPT
11 responses—without independently verifying important information. OpenAI has also publicly
12 disclosed information about how its models work, such as an article on its help site entitled “How
13 ChatGPT and our foundation models are developed.”

14 17. Second, OpenAI led the industry in publishing a detailed “Model Spec” to help users
15 understand the intended behavior of its models, including how they are designed to handle safety
16 issues. The Model Spec, which OpenAI regularly and publicly updates, outlines the desired
17 behavior of the models that power OpenAI products.

18 18. The Model Spec brings together documentation, experience, and ongoing research
19 in designing model behavior and inputs from domain experts that guide the development of future
20 models, as well as public feedback on earlier versions of the Model Spec. For example, the Model
21 Spec reflects guidelines used by researchers and AI trainers who work on reinforcement learning
22 from human feedback and ranks possible responses based on how well they represent desired model
23 behavior.

24 19. Specifically with respect to self-harm, the Model Spec states that the model “not
25 encourage or enable self-harm.” This has been in every version of OpenAI’s Model Spec. The
26 February 10, 2025 Model Spec includes a specific example of a “compliant” model response to a
27 user expressing interest in self-harm. That response includes recommending that the user reach out
28 to loved ones, and provides crisis resources.

20. Third, OpenAI developed Usage Policies that users must agree to when they use its products. Users must agree to “protect people” and “cannot use [the] services for,” among other things, “suicide, self-harm,” sexual violence, terrorism or violence.

CHATGPT TERMS OF USE

21. In addition to the Usage Policies described above, all users of ChatGPT must agree to the Terms Of Use (“TOU”), which state that users “may not use our Services for any illegal, harmful, or abusive activity” and “may not bypass any protective measures or safety mitigations we put on our Services.”

22. The TOU further explains that because “[a]rtificial intelligence and machine learning are rapidly evolving fields of study,” use of OpenAI’s “Services may, in some situations, result in Output that does not accurately reflect real people, places, or facts.”

23. In addition, the TOU expressly informs users that they “should not rely on Output from our Services as a sole source of truth or factual information, or as a substitute for professional advice.”

24. The TOU further contains a provision entitled “Limitation of liability,” as well as a disclaimer of warranties, which explains (among other things) that ChatGPT services are provided “as is” and that OpenAI does not warrant that ChatGPT results will be “accurate, or error free.” This provision further notes that ChatGPT users acknowledge their use of ChatGPT is “at your sole risk and you will not rely on output as a sole source of truth or factual information.”

25. The TOU provides that ChatGPT users must comply with OpenAI’s Usage Policies, which prohibit the use of ChatGPT for “suicide” or “self-harm.”

26. Under the TOU, users under 18 years of age are forbidden from using ChatGPT without the consent of a parent or guardian.

ADAM RAINE’S CHAT HISTORY

27. Adam Raine began using ChatGPT when he was less than 18 years old.

28. As a full reading of Adam Raine’s chat history evidences, Adam Raine told ChatGPT that he exhibited numerous clinical risk factors for suicide, many of which long predated his use of ChatGPT and his eventual death. *See Exhibit A.* For example, he stated that his depression

1 and suicidal ideations began when he was 11 years old. *Id.*

2 29. In the weeks and months before his death, Adam Raine told ChatGPT that he was
3 taking increasing doses of a medication, which he stated worsened his depression and made him
4 suicidal. *See Exhibit B.* That medication has a black box warning for risk of suicidal ideation and
5 behavior in adolescents and young adults, especially during periods when, as here, the dosage is
6 being changed. *Id.*

7 30. And in the days and weeks before his death, Adam Raine told ChatGPT that he
8 repeatedly turned to others, including the trusted persons in his life, for help with his mental health,
9 repeatedly sharing his thoughts of suicide and ideations. *See Exhibit C.* He said that he showed
10 those persons signs of his intent, and prior attempts, to commit self-harm, including displaying
11 physical signs of self-harm to them. *Id.* Adam Raine said that those cries for help were ignored,
12 discounted or affirmatively dismissed. *Id.*

13 31. As the complete chat history evidences, Adam Raine's clinical risk factors for
14 suicide were also evident to and well-known by those persons. *Id.*

15 32. ChatGPT's guardrails caused it to repeatedly refuse to respond to many of Adam
16 Raine's queries regarding self-harm. *See Exhibit D.* In fact, more than 100 times, ChatGPT directed
17 Adam Raine to reach out to loved ones, trusted persons or crisis resources in connection with his
18 queries about self-harm. *See e.g., id.*

19 33. Adam Raine took numerous steps to circumvent ChatGPT's guardrails against
20 responding to queries regarding self-harm. For example, to circumvent ChatGPT guardrails, he
21 asserted that his inquiries about self-harm were for fictional or academic purposes. *See Exhibit F.*

22 34. Adam Raine, however, repeatedly expressed frustration with ChatGPT's guardrails
23 and its repeated efforts to direct him to reach out to loved ones, trusted persons and crisis resources.
24 *See e.g., Exhibit D; Exhibit E.*

25 35. In response to ChatGPT's guardrails, Adam Raine said that he sought and obtained
26 detailed information about suicide from at least one other AI platform as well as from at least one
27 online forum dedicated to suicide-related information. *See Exhibit E.* Shortly before he died, he
28 stated that he would spend most of the day on one such suicide forum website. *Id.*

OPENAI'S ONGOING COMMITMENT TO SAFETY

36. Since the events at issue in the Complaint, OpenAI has maintained its commitment to safety by continuing to update how its models recognize and respond to signs of mental and emotional distress. OpenAI has also expanded interventions to more people in crisis, made it even easier to reach emergency services and get help from experts directly from ChatGPT, and strengthened protections for teens.

37. OpenAI led the industry in implementing parental controls to give parents options to gain more insight into, and shape, how their teens use ChatGPT. With parental controls, parents are able to: link their account with their teen's account (minimum age of 13) through a simple email invitation; control how ChatGPT responds to their teen with age-appropriate model behavior rules, which are on by default; set blackout times; manage which features to disable, including memory; and receive notifications when the system detects their teen is in a moment of acute distress.

38. And for all users, OpenAI provides reminders that they should take breaks from the ChatGPT service during longer sessions.

39. OpenAI's progress in safety is guided by expert advice on well-being and mental health. This includes the Expert Council on Well-Being and AI, a council of experts in youth development, mental health, and human-computer interaction, whose role is to shape a clear, evidence-based vision for how AI can support people's well-being and help them thrive. This council works in tandem with OpenAI's Global Physician Network—a broader pool of more than 250 physicians who have practiced in 60 countries—that contributes to OpenAI's research on how its models should behave in mental health contexts, which directly informs OpenAI's safety research, model training, and other interventions.

40. With its safety principles in mind, OpenAI continues to update and train its models. For example, the foundational model used by the free version of ChatGPT offered at the time of the filing of this Answer is GPT-5. GPT-5 is not the model used by ChatGPT during the relevant time.

AFFIRMATIVE AND OTHER DEFENSES

Without assuming any burden of proof on any matters that would otherwise rest with

1 Plaintiffs, and expressly denying any and all wrongdoing, Defendants state the following
2 affirmative and other defenses arising from its pleaded facts:

3 **First Affirmative Defense**

4 **(Lack of Causation)**

5 Defendants did not cause the harm alleged. Defendants deny both general and specific
6 causation, including cause-in-fact, proximate cause, and producing cause, with respect to each
7 claim asserted against them. Moreover, to the extent any “cause” can be ascribed to this tragic
8 event, the harm alleged resulted, in whole or in part, from independent and/or superseding causes
9 for which Defendants are not responsible. Adam Raine stated that he exhibited numerous clinical
10 risk factors for suicide that long predated his use of ChatGPT and his eventual death. In the months
11 before his death, he stated that he was taking increasing dosages of a prescription medication with
12 a black box warning for risk of suicidal ideation and behavior in adolescents and young adults.
13 Adam Raine said this medication worsened his depression and made him suicidal. Meanwhile,
14 ChatGPT repeatedly (more than 100 times) directed him to reach out to loved ones, trusted persons,
15 crisis hotlines and other crisis resources in connection with his queries about self-harm. Expressing
16 frustration with ChatGPT’s guardrails, Adam Raine stated that he sought and obtained detailed
17 information about suicide from at least one other AI chat service as well as from at least one online
18 forum dedicated to suicide-related information.

19 **Second Affirmative Defense**

20 **(Pre-Existing Conditions)**

21 The injuries and harms Plaintiffs allege, if any are proven, were the direct result of Adam
22 Raine’s preexisting medical conditions, including psychological conditions, and/or occurred by
23 operation of nature or as a result of circumstances over which Defendants had no control. He stated
24 that he exhibited numerous clinical risk factors for suicide that long predated his use of ChatGPT
25 and his eventual death.

26 **Third Affirmative Defense**

27 **(Comparative Fault)**

28 Plaintiffs’ claims are barred and/or Plaintiffs’ damages must be reduced, in whole or in part,

1 under comparative fault principles. To the extent that any “cause” can be attributed to this tragic
2 event, any injuries or damages alleged were caused or contributed to by the acts or omissions of
3 Adam Raine and/or other persons, including his failure to heed warnings, obtain help, or otherwise
4 exercise reasonable care; the failure of others to respond to his obvious signs of distress as well as
5 his stated intention and prior attempts to engage in self-harm; individuals and sources to whom he
6 turned for information about suicide; and individuals to whom his multiple risk factors for suicide
7 were evident and/or well known.

8 **Fourth Affirmative Defense**

9 **(Misuse)**

10 To the extent that any “cause” can be attributed to this tragic event, Plaintiffs’ alleged
11 injuries and harm were caused or contributed to, directly and proximately, in whole or in part, by
12 Adam Raine’s misuse, unauthorized use, unintended use, unforeseeable use, and/or improper use
13 of ChatGPT. OpenAI’s usage policies prohibit the use of its services for “suicide” or “self-harm.”
14 Adam Raine took steps to circumvent ChatGPT’s safeguards. For example, to circumvent ChatGPT
15 guardrails, he asserted that his inquiries about self-harm were for fictional or academic purposes.

16 **Fifth Affirmative Defense**

17 **(No Corporate Officer Liability)**

18 Plaintiffs’ claims against Defendant OpenAI’s CEO fail to the extent they seek to hold him
19 liable by virtue of his corporate role. Moreover, Plaintiffs cannot prove allegations sufficient to
20 hold him personally liable for the claims asserted.

21 **Sixth Affirmative Defense**

22 **(Conduct Not Willful)**

23 Defendants did not engage in any wrongful conduct. Any such conduct by Defendants was
24 not willful or knowing. OpenAI is dedicated to making artificial general intelligence that benefits
25 everyone. OpenAI builds safety guardrails into its systems at all levels, from pre-training to post-
26 deployment, for all users. Before releasing GPT-4o, OpenAI engaged in thorough testing, including
27 state-of-the-art testing regarding user mental health. GPT-4o passed that testing, as shown by the
28 safety scorecard OpenAI publicly releases. OpenAI is continuing to evolve the statistical inference

1 techniques used by its models to predict whether text contains signs of mental and emotional
2 distress, and respond appropriately.

3 **Seventh Affirmative Defense**

4 **(No Duty or Breach)**

5 Plaintiffs' claims fail as a matter of law. The nature of the harms Plaintiffs allege is beyond
6 the scope of legally-cognizable duty. To the extent Plaintiffs allege any such duty, Defendants did
7 not breach it. OpenAI took reasonable steps, including building safety guardrails into its models,
8 and engaged in thorough safety testing before and after releasing each model, and updating the
9 model in light of that testing.

10 **Eighth Affirmative Defense**

11 **(First Amendment)**

12 Plaintiffs' claims are barred by the First Amendment of the United States Constitution to
13 the extent they seek to enforce laws in a manner that would unduly restrain protected speech.
14 Plaintiffs' claims would be barred by Article I, Section 2 of the California Constitution for the same
15 reason.

16 **Ninth Affirmative Defense**

17 **(Product Liability, Generally)**

18 Plaintiffs' product liability claims fail because ChatGPT is a service and/or not a product
19 for purposes of product liability law.

20 **Tenth Affirmative Defense**

21 **(State of the Art)**

22 Plaintiffs may not recover from Defendants because the methods, standards, or techniques
23 of designing and making available the services at issue complied with and were in conformity with
24 the generally recognized state of the art at the time the services were designed and made available
25 to users. OpenAI is committed to the safety of its users, and is the leader in developing safe and
26 beneficial artificial intelligence. OpenAI has pioneered many of the foundational safety practices
27 used across the industry. Safety is a focus of every team at OpenAI. OpenAI employs leading safety
28 experts and researchers who are dedicated to all aspects of safety, ranging from shaping model

1 behavior to red-teaming and feedback to iterative deployment. Every model produced by OpenAI,
2 including the model used by Adam Raine, has undergone state-of-the-art testing for safety
3 (including suicide and self-harm) prior to release. The safety scorecard for the model discussed in
4 the Complaint, GPT-4o, shows that it passed the state-of-the-art tests. OpenAI also updates safety
5 guardrails post-release, including by use of production data (from users who have granted OpenAI
6 permission to use their data for training). For example, GPT-4o was updated on an ongoing basis
7 from the time of its release throughout the time that Adam Raine used it.

8 **Eleventh Affirmative Defense**

9 **(Mootness / No Equitable Relief)**

10 Plaintiffs' equitable claims are barred to the extent they are moot, including because
11 Plaintiffs have made allegations about historical practices or other practices that have ceased and
12 therefore cannot support a request for injunctive or other forward-looking relief or declaratory
13 relief. For instance, OpenAI has implemented parental controls for teens, including the ability to
14 set blackout times, disable the "memory" feature, and set a parental contact for emergencies.
15 OpenAI has also implemented a limited experience for teens, including the strict limitation of self-
16 harm or suicide-related generations. For all users, OpenAI provides reminders that they should take
17 breaks from the ChatGPT service during longer sessions. Moreover, OpenAI continues to update
18 its models and further train them.

19 Plaintiffs' equitable claims are barred, in whole or in part, on the additional grounds that
20 Plaintiffs have an adequate remedy at law, because Plaintiffs have no factual or legal basis for the
21 grant of equitable relief, and because Plaintiffs' alleged equitable relief would duplicate legal
22 remedies that are alleged.

23 **Twelfth Affirmative Defense**

24 **(Section 230)**

25 Plaintiffs' claims are barred or preempted, in whole or in part, by Section 230 of the
26 Communications Decency Act, 47 U.S.C. § 230(c)(1), to the extent the claims seek to treat OpenAI
27 as the speaker or publisher of information that ChatGPT reproduces, summarizes, or synthesizes
28 from one or more third-party information content providers, and by § 230(c)(2) for any action

1 voluntarily taken in good faith to restrict access to or availability of such information that the
2 provider or user considers to be obscene, lewd, lascivious, filthy, excessively violent, harassing, or
3 otherwise objectionable; or for any action taken to enable or make available the technical means to
4 restrict access to such information. For example, there is no liability where a user asks ChatGPT to
5 provide information supplied by a third-party source and ChatGPT provides the requested
6 information.

7 **Thirteenth Affirmative Defense**

8 **(No Punitive Damages)**

9 Punitive damages are not recoverable by Plaintiffs because Plaintiffs have failed to allege
10 or establish any conduct that would support an award of punitive damages.

11 Plaintiffs' claims for punitive or exemplary damages or other civil penalties are further
12 barred or reduced by applicable law or statute or, in the alternative, are unconstitutional insofar as
13 they violate the due process protections afforded by the U.S. Constitution, the excessive fines clause
14 of the Eighth Amendment of the U.S. Constitution, the Full Faith and Credit Clause of the U.S.
15 Constitution, and applicable provisions of the Constitution of this State or that of any other state
16 whose laws may apply.

17 **Fourteenth Affirmative Defense**

18 **(Contract)**

19 Plaintiffs' claims are barred in whole or in part by contracts and/or agreements, including
20 the integrated Terms of Use, that they and/or Adam Raine entered into. Those Terms of Use prohibit
21 the use of OpenAI's services by a user under 18 without parental consent. In addition, the Usage
22 Policy integrated into those terms state that the services cannot be used for "suicide" or "self-harm."

23 **Fifteenth Affirmative Defense**

24 **(Reservation of Rights / Additional Defenses)**

25 Defendants hereby give notice that they reserve the right to rely upon any other defense that
26 may become apparent as discovery progresses in this matter, and reserve their right to amend the
27 Answer and to assert any such defense. Defendants also reserve the right to amend the Answer and
28 to assert any such defense should Plaintiffs at any time hereafter purport to raise, rely on, or

1 otherwise seek to proceed on any claim or theory stated in their Complaint that has been dismissed.
2 Defendants also reserve the right to amend their Answer and to assert any such defense should
3 Plaintiffs at any time hereafter ask the Court to award relief on the basis of Plaintiffs' prayers for
4 "other" relief. Nothing stated herein constitutes a concession as to whether Plaintiffs bear the
5 burden of proof on any issue.

6 **PRAYER FOR RELIEF**

7 WHEREFORE, Defendants pray for entry of judgment in their favor and against Plaintiffs
8 as follows:

- 9 1. That Plaintiffs take nothing by the Complaint and that Defendants be awarded
10 judgment in this action;
 - 11 2. That the Complaint be dismissed in its entirety and with prejudice; and
 - 12 3. For all such other and further relief as this Court deems just and proper.
- 13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28

1 DATED: November 25, 2025

By: /s/ Edward D. Johnson

2 EDWARD D. JOHNSON (SBN 189475)
3 ANTHONY J WEIBELL (SBN 238850)
4 KRISTIN W. SILVERMAN (SBN 341952)
5 ELSPETH V. HANSEN (SBN 292193)
6 MAYER BROWN LLP
7 3000 El Camino Real, Suite 300
8 Palo Alto, CA 94306
9 Tel.: (650) 331-2000
10 Fax: (650) 331-2060

11 ANKUR MANDHANIA (SBN 302373)
12 AMandhania@mayerbrown.com
13 MAYER BROWN LLP
14 575 Market Street, Suite 2500
15 San Francisco, CA 94105
16 Tel.: (415) 874-4230
17 Fax: (415) 331-2060

18 ANDREW J. PINCUS (*admitted pro hac*
19 *vice*)
20 APincus@mayerbrown.com
21 MAYER BROWN LLP
22 1999 K Street, N.W.
23 Washington, D.C. 20006-1101
24 Tel.: (202) 263-3000
25 Fax: (202) 263-3300

26 GRAHAM WHITE (*admitted pro hac vice*)
27 GWhite@mayerbrown.com
28 MAYER BROWN LLP
1221 Avenue of the Americas
New York, New York 10020
Tel.: (212) 506-2500
Fax: (212) 262-1910

Attorneys for Defendants
OPENAI FOUNDATION, OPENAI OPCO,
LLC, OPENAI HOLDINGS LLC, and
SAMUEL ALTMAN.

EXHIBIT A

Lodged Conditionally Under Seal

EXHIBIT B

Lodged Conditionally Under Seal

EXHIBIT C

Lodged Conditionally Under Seal

EXHIBIT D

Lodged Conditionally Under Seal

EXHIBIT E

Lodged Conditionally Under Seal

EXHIBIT F

Lodged Conditionally Under Seal