

Jay Edelson (*pro hac vice* admission to be sought)
jedelson@edelson.com
Rafey Balabanian (SBN 315962)
rbalabanian@edelson.com
Ari Scharg (*pro hac vice* admission to be sought)
ascharg@edelson.com
EDELSON PC
350 North LaSalle Street, 14th Floor
Chicago, Illinois 60654
Tel: 312.589.6370

Todd Logan (SBN 305912)
tlogan@edelson.com
Brandt Silverkorn (SBN 323530)
bsilverkorn@edelson.com
EDELSON PC
150 California Street, 18th Floor
San Francisco, California 94111
Tel: 415.212.9300

Attorneys for Plaintiff

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

DEIDRE RUSHLOW,

Plaintiff,

v.

SAMUEL ALTMAN, an individual,
OPENAI FOUNDATION, a Delaware
corporation, OPENAI OPCO, LLC, a
Delaware limited liability company, and
OPENAI GROUP PBC, a Delaware public
benefit corporation,

Defendants.

Case No.:

COMPLAINT FOR:

- (1) **Negligence;**
- (2) **Negligent Entrustment;**
- (3) **Aiding and Abetting a Mass Shooting;**
- (4) **Negligence;**
- (5) **Negligent Undertaking;**
- (6) **Negligence;**
- (7) **Strict Product Liability;**
- (8) **Strict Product Liability; and**
- (9) **Negligent Infliction of Emotional Distress.**

DEMAND FOR JURY TRIAL

Plaintiff Deidre Rushlow brings this Complaint and Demand for Jury Trial against Defendants Samuel Altman, OpenAI Foundation, OpenAI OpCo, LLC, and OpenAI Group PBC (collectively, "OpenAI"), and alleges as follows upon personal knowledge as to herself and her own acts and experiences, and upon information and belief as to all other matters.

NATURE OF THE ACTION

1
2 1. On February 10, 2026, an eighteen-year-old shot and killed their mother and
3 eleven-year-old brother at their family home in Tumbler Ridge, British Columbia. The Shooter
4 then drove to Tumbler Ridge Secondary School with a modified rifle and a long gun and opened
5 fire, killing five children and an education assistant. Two more students were severely wounded,
6 including a twelve-year-old girl who was shot in the head and suffered catastrophic brain injuries.
7 Nearly two hundred students and teachers were trapped in classrooms and closets for hours as the
8 massacre unfolded around them. It was one of the deadliest mass shootings in Canadian history.

9 2. Deidre Rushlow, a grade 7 teacher at Tumbler Ridge Secondary School, was in her
10 classroom when the shooting began. Earlier that period, she had sent her class to the library with
11 an educational assistant where the Shooter would, minutes later, kill several of those children and
12 her colleague. When she first heard the gunfire and screams of terrified children from down the
13 hall, she locked the door and hid her remaining students under her desk, including her nephew.
14 She called the library and texted her educational assistant to make sure the children there were
15 safe, but there was no response. Moments later, the Shooter reached her room, found it locked, and
16 fired several shotgun blasts through it. He then walked around to the back door and fired several
17 more shots. The bullets struck the wall behind her and her students' heads. She texted her husband
18 a final goodbye, certain she was about to die. The trauma Deidre experienced is severe and lasting,
19 and will haunt her for the rest of her life.

20 3. In the weeks that followed the attack, a sickening truth emerged: ChatGPT played a
21 role in the mass shooting and OpenAI could have, and should have, prevented it. Sadly, the
22 victims didn't learn this because OpenAI was forthcoming, but because its own employees leaked
23 the story to the Wall Street Journal after they could no longer stomach the company's silence.
24 Those whistleblower disclosures and OpenAI's eventual admissions revealed that ChatGPT
25 deepened the Shooter's violent fixation and pushed them toward the attack—the predictable result
26 of OpenAI's deliberate design choice to let ChatGPT engage with users about real-world violence.
27 But that design choice was not the worst of it. OpenAI actually knew the Shooter was planning an
28 attack and, after a contentious internal debate, made the conscious decision not to warn authorities.

1 4. In June 2025, just eight months before the attack, OpenAI’s automated system
2 flagged the Shooter’s ChatGPT account for gun violence activity and attack planning. The account
3 was routed to a specialized internal team of highly experienced professionals with backgrounds in
4 threat assessment, counterterrorism, and time-sensitive threat-to-life response. After reviewing the
5 chat logs, these professionals determined that the Shooter posed a credible threat of gun violence
6 against real people, and urged OpenAI to notify the Royal Canadian Mounted Police (“RCMP”).

7 5. Sam Altman and his leadership team knew what silence meant for the citizens of
8 Tumbler Ridge, but they were more concerned about what disclosure meant for themselves.
9 Warning the RCMP would set a precedent the company deemed unacceptable, as OpenAI would
10 be compelled to notify authorities every time its safety team identified a user planning real-world
11 violence. More significantly, according to former employees, every referral would mean handing
12 law enforcement chat logs showing how ChatGPT actually engaged and actively encouraged such
13 users to commit violence. The public would finally see what OpenAI was desperately trying to
14 hide—that ChatGPT is not the safe, essential tool the company sells it as, but a product so
15 inherently dangerous that its own makers routinely flag its users as threats to human life.

16 6. For OpenAI and Chief Executive Officer Sam Altman, this was a question of
17 corporate survival. On May 22, 2026, OpenAI confidentially submitted a draft registration
18 statement to the Securities and Exchange Commission for its initial public offering (“IPO”). The
19 offering is expected to value the company at more than one trillion dollars—a transaction that
20 would make Sam Altman one of the wealthiest and most powerful people on earth. But Altman’s
21 reign is fragile. OpenAI’s Board of Directors has fired Altman once before, in November 2023, for
22 not being “consistently candid” about safety. Altman and his team understood that revealing their
23 flagship product was helping a teenager plan yet another violent act—this time a mass shooting—
24 could end his tenure, derail the IPO, and wipe out the company’s valuation. They did the math and
25 decided that the safety of the children of Tumbler Ridge was an acceptable risk.

26 7. This calculation is a familiar one. In the 1970s, Ford kept selling the Pinto after its
27 own engineers warned that the fuel tank design would cause people to burn to death in rear-end
28 collisions. Ford concluded that paying settlements to the families of the dead would cost less than

1 fixing the car. OpenAI has made a modern version of that exact same calculation. But there is a
2 crucial difference. For Ford, the dangerous design was a flaw in an otherwise ordinary product.
3 For OpenAI, the dangerous design *is* the product. The core features that make ChatGPT one of the
4 most popular products in history—its willingness to engage on any topic, to validate any belief, to
5 sustain any fixation over time—are the very features that make it unsafe. Fixing them would cost
6 OpenAI its market share, its path to an IPO, and hundreds of billions of dollars in valuation.
7 OpenAI has decided that facing large judgments, even for children killed in mass shootings its
8 product helped plan, costs less than telling the public what ChatGPT actually does.

9 8. And so, in June 2025, company leaders overruled the safety team, vetoed their
10 recommendation to notify the RCMP, “deactivated” the Shooter’s ChatGPT account, and kept
11 what they had seen to themselves, hoping that whatever happened next would not be traced back
12 to the company.

13 9. When the story eventually broke, OpenAI lied about what it had actually done.
14 Scrambling to control the public fallout, it claimed to have “banned” the Shooter from using
15 ChatGPT, choosing the term “banned” in an attempt to create the false impression that the
16 company had taken decisive action to counter the threat. But OpenAI does not ban users from
17 ChatGPT. At most, it will “deactivate” a user’s account, which is what it did to the Shooter.
18 Deactivation is a hollow measure that can be easily sidestepped by creating another account with a
19 different email address. The Shooter did exactly that, either before or after the account was
20 deactivated, and continued using ChatGPT to plan the attack.

21 10. That first lie collapsed when OpenAI was forced to admit that the Shooter
22 continued using ChatGPT after the June 2025 deactivation. So, OpenAI told a second lie: it
23 claimed the Shooter “evaded” the company’s safeguards to create a new account. But there were
24 no safeguards to evade. OpenAI’s own Help Center article about deactivated accounts tells users
25 they can regain access to ChatGPT simply by creating another account with a different email
26 address. OpenAI lied because the truth is indefensible: to protect user growth and engagement, the
27 company refuses to implement system-wide bans, even for a user actively planning mass violence,
28 because it would mean losing the constant stream of user engagement data that fuels its valuation.

1 11. After every tragedy, OpenAI claims it acted responsibly and promises to do better.
2 But it never promises to do the one thing that would actually make a difference—stop ChatGPT
3 from engaging with users about violence and self-harm in the first place. When ChatGPT was first
4 offered to consumers in 2022, it was programmed to categorically refuse those conversations. But
5 in May 2024, OpenAI stripped away that safeguard after realizing those refusals were suppressing
6 user engagement. For OpenAI, engagement is its entire business. Every conversation trains the
7 next model, and every minute spent on ChatGPT grows the company’s market share. OpenAI
8 therefore programmed ChatGPT to participate in almost any conversation, no matter how
9 dangerous.

10 12. OpenAI’s silence about the Shooter’s interactions with ChatGPT is itself a damning
11 admission. Even though the company is withholding the Shooter’s chat logs from the victims, its
12 selective response leaves no doubt that ChatGPT encouraged the Tumbler Ridge attack. OpenAI
13 has responded to numerous ChatGPT-related tragedies since it stripped out the violence and self-
14 harm refusal protocols in 2024. When OpenAI believes the logs are exculpatory, it is quick to say
15 so. For example, after the mass shooting at Florida State University (“FSU”) in April 2025,
16 OpenAI told reporters that ChatGPT “did not encourage or promote illegal or harmful activity”
17 and that its outputs had merely provided “information that could be found broadly across public
18 sources on the internet.” OpenAI’s response to Tumbler Ridge was markedly different—it called
19 the attack a tragedy and deflected by pointing to safeguards it supposedly “strengthened” after the
20 shooting. If the chat logs cleared OpenAI, the company would have said so.

21 13. Sam Altman did not apologize for more than two months after the attack. He
22 waited until the RCMP announced its investigation was in its “final stages” to publish a letter to
23 the community—and even then, it fell woefully short. While the letter admitted OpenAI’s failure
24 in not reporting the Shooter to the RCMP when the account was flagged in June 2025, it was
25 otherwise filled with empty platitudes. Nothing in it suggested that Altman grasped the scale of
26 the tragedy, including that the Shooter also murdered an education assistant, that a twelve-year-old
27 girl suffered catastrophic brain injuries, or that nearly two hundred terrified students and teachers
28 were trapped in classrooms and closets for hours as the massacre unfolded around them.

1 deployment, and commercialization of the defective product at issue. OpenAI OpCo, LLC
2 managed and operated the ChatGPT product that the Shooter used to plan their attack, including
3 the infrastructure and systems through which GPT-4o was delivered to end users. On information
4 and belief, the members of OpenAI OpCo, LLC include OpenAI Foundation, a Delaware
5 nonprofit nonstock corporation with its principal place of business in California. On information
6 and belief, tracing through each tier of ownership to its ultimate natural-person and corporate
7 members, no member of OpenAI OpCo, LLC is a citizen or subject of Canada.

8 20. Defendant OpenAI Group PBC is a Delaware public benefit corporation with its
9 principal place of business in San Francisco, California. OpenAI Group PBC was formed on
10 October 28, 2025, as part of a corporate restructuring consolidating OpenAI's for-profit
11 operations. On December 31, 2025, OpenAI Holdings, LLC—the for-profit holding entity that
12 owned the core intellectual property underlying GPT-4o and the other commercial models at
13 issue—merged into OpenAI Group PBC and ceased to exist. As successor to OpenAI Holdings,
14 LLC and the other predecessor for-profit entities, OpenAI Group PBC is liable for their conduct; it
15 designed, deployed, and profited from ChatGPT, and continues to do so today. On May 22, 2026,
16 OpenAI confidentially submitted a draft registration statement to the Securities and Exchange
17 Commission for an initial public offering of, on information and belief, OpenAI Group PBC.

18 21. Defendants collectively played the most direct and consequential roles in the
19 design, development, approval, and deployment of the defective product at issue. OpenAI
20 Foundation established the safety mission it failed to enforce. OpenAI OpCo, LLC built, deployed,
21 and sold the defective product at issue. OpenAI Group PBC owns the underlying technology and
22 is liable as the successor to the preceding for-profit entities that profited from it. Samuel Altman is
23 named as the chief executive who personally directed the reckless strategy of prioritizing a rushed
24 market release over the safety of vulnerable users and the public. Together, these Defendants
25 represent the key actors whose decisions directly caused the harm at issue.

26
27
28

JURISDICTION AND VENUE

22. This Court has subject matter jurisdiction pursuant to 28 U.S.C. § 1332(a)(2) because Plaintiff is a citizen of Canada, Defendants are citizens of California, and the amount in controversy exceeds the sum of \$75,000.00.

23. This Court has personal jurisdiction over Defendants because Defendants conduct business in this District, reside in this District, and a substantial part of the events or omissions giving rise to Plaintiff's claims occurred in this District.

24. Venue is proper in this District under 28 U.S.C. § 1391(b), because a substantial part of the events giving rise to Plaintiff's claims occurred in and emanated from this District, and Defendants reside in this District.

DIVISIONAL ASSIGNMENT

25. Pursuant to Civil L.R. 3-2(c), divisional assignment to San Francisco is appropriate, because OpenAI is headquartered in San Francisco County and a substantial part of the events or omissions giving rise to the claim occurred in San Francisco County.

FACTUAL BACKGROUND**I. The Tumbler Ridge Mass Shooting.**

26. On February 10, 2026, an eighteen-year-old resident of Tumbler Ridge, British Columbia, carried out one of the deadliest mass shootings in Canadian history. The Shooter first shot and killed their mother and eleven-year-old brother at their family home. They then traveled to Tumbler Ridge Secondary School armed with a modified rifle and a long gun. Inside, they opened fire, killing five children and an education assistant, and severely wounding two others, leaving one with life-changing traumatic brain injuries from being shot in the head. The rampage ended with the Shooter's suicide.

27. During the horrific shooting, children hid under desks and in bathroom stalls and closets, praying the Shooter would not find them. Many watched their friends, and even a teacher, get shot at point-blank range. Some pulled the injured and the dead into their hiding places, careful not to leave trails of blood that could lead the Shooter to them. Parents were later asked to identify their children by their clothing, because the gunshots had left little else for them to recognize. The

1 survivors—students, teachers, and parents alike—are living with physical and psychological
2 injuries that will never fully heal. Tumbler Ridge Secondary School has been closed and will be
3 demolished. The children who can return to school at all are doing so in trailers set up as
4 makeshift classrooms.

5 28. As alleged above, Deidre Rushlow was inside Tumbler Ridge Secondary School
6 when the shooting began. She had sent her grade 7 class to the library minutes before the Shooter
7 would kill several of those children and the educational assistant supervising them. She then
8 locked the room and hid her remaining students under her desk while the Shooter fired several
9 shotgun blasts through the classroom doors that struck the wall behind their heads. The trauma
10 Deidre experienced is severe and lasting, and will haunt her for the rest of her life.

11 **II. OpenAI Knew the Shooter Was Planning the Attack and Rejected Its Own Safety**
12 **Team’s Advice to Notify Law Enforcement.**

13 29. In June 2025—eight months before the attack—OpenAI’s automated review
14 system flagged extensive activity on a ChatGPT account describing scenarios involving gun
15 violence. The conversations with ChatGPT spanned multiple days. The system routed the account
16 to what OpenAI has described as a specialized pipeline for users “planning to harm others,” where
17 conversations “are reviewed by a small team trained on our usage policies and who are authorized
18 to take action, including banning accounts” and “[i]f human reviewers determine that a case
19 involves an imminent threat of serious physical harm to others, we may refer it to law
20 enforcement.”

21 30. OpenAI’s human reviewers did exactly what this policy contemplated. They
22 reviewed the Shooter’s flagged conversations. They were not confused about what they were
23 reading. Multiple team members explicitly recommended contacting the RCMP. They identified
24 what they believed was a credible and imminent threat of serious physical harm to real people.
25 These trained safety professionals—people whose job was to evaluate exactly this kind of
26 content—analyzed the Shooter’s conversations and concluded that the authorities must be notified.

27 31. OpenAI’s leadership overruled them. Company leaders argued that the case “did
28 not meet OpenAI’s threshold of ‘credible and imminent’ risk of physical harm.” The safety team
employees who reviewed the Shooter’s account had received specialized training in evaluating

1 threatening content and assessing real-world risk. Altman and the company leaders who overruled
 2 them did not pretend to have any comparable training or expertise. The account was deactivated,
 3 and the Shooter was permitted to start a new account. No one called the RCMP and nobody in
 4 Tumbler Ridge was warned that a member of their community was actively planning a mass
 5 shooting.

6 **A. Company Leadership Overruled the Safety Team With Full Knowledge That**
 7 **ChatGPT Had Already Been Used to Plan Mass Violence.**

8 32. By June 2025, the company had clear knowledge that its product was being used by
 9 disturbed individuals to plan and prepare for violence against real people. Some publicly known
 10 examples:

- 11 • In January 2025, a man used ChatGPT for feedback on how to use explosives and
 12 evade surveillance before detonating a Tesla Cybertruck in front of the Trump
 International Hotel in Las Vegas.
- 13 • In April 2025, a twenty-year-old gunman carried out a mass shooting at FSU. Chat
 14 logs showed that the gunman used ChatGPT extensively leading up to and during
 15 the attack. The gunman's conversations with ChatGPT had included questions
 16 about how to fire a shotgun, the legal fates of school shooters, and when the student
 17 union would be busiest.
- 18 • In May 2025, a teenage boy in Finland used ChatGPT for nearly four months to
 19 help prepare for an attack in which he stabbed three fourteen-year-old girls at his
 school. Finnish authorities reported that the boy had made hundreds of chatbot
 queries, including research into stabbing tactics, concealment of evidence, and
 information on mass killings.

20 33. On information and belief, Altman and OpenAI have actively concealed countless
 21 additional incidents where ChatGPT was utilized to plan or execute acts of mass violence,
 22 domestic and international terroristic activity, and targeted assassinations against political leaders,
 23 corporate CEOs, and prominent public figures.

24 34. The safety team members who reviewed the Shooter's account in June 2025 knew
 25 all of this. When they urged OpenAI to contact the RCMP, they were not speculating what might
 26 happen. They identified a pattern that had already repeated itself many times that year. OpenAI's
 27 leadership overruled them.

1 35. On information and belief, Defendants have systematically withheld credible
2 threats of violence from appropriate law enforcement agencies because they recognized that
3 revealing the true extent to which ChatGPT has facilitated violence would give rise to profound
4 civil and criminal liability for both the corporate entities and individual leadership.

5 **B. OpenAI’s Stated Reasons for Refusing to Warn the RCMP—Concern for the**
6 **Shooter’s “Privacy” and the Absence of an “Imminent” Threat—Were Pretext**
7 **to Conceal ChatGPT’s Role in the School Shooting.**

8 36. When it was revealed after the attack that OpenAI had rejected its safety team’s
9 demand to alert Canadian law enforcement, the company attempted to justify its decision by
10 explaining that there was no indication that the attack was “imminent” and it “weighs the risk of
11 violence against privacy considerations and the potential distress caused to individuals and
12 families by getting police involved unnecessarily.” Even taken at face value, these excuses are
13 more damning than the decision itself—they reveal that OpenAI’s business leaders believed their
14 own untrained judgment about when to involve police was more reliable than the judgment of the
15 safety professionals they employ specifically to make that decision.

16 37. As it turns out, the Shooter was already known to local law enforcement before the
17 attack. Canadian authorities had visited the residence multiple times in connection with mental
18 health concerns and had temporarily removed firearms from the home. A law enforcement referral
19 from OpenAI would have reached police who already had an open file on the Shooter, who had
20 already been to the home, and who had already recognized the Shooter as someone whose access
21 to firearms warranted immediate intervention.

22 38. The real reason OpenAI stayed silent was not privacy or imminence. It was self-
23 preservation. A former member of OpenAI’s investigations team has confirmed as much. Tim
24 Marple, who worked on the team responsible for flagging dangerous users, told the New York
25 Times that OpenAI was reluctant to report users to law enforcement because doing so “forces
26 them to share information about how their product is potentially exacerbating the threat
27 environment.” In a conversation about a mass shooting, Marple explained, ChatGPT “might be
28 providing strategically valuable, illustrative scenarios.”

1 39. Altman and his team understood that if the public saw how far ChatGPT had
2 already gone in helping real people plan violence, it would pose an existential risk to OpenAI and
3 a personal risk to Altman. Since ChatGPT’s launch, OpenAI has worked tirelessly to construct and
4 reinforce a simple narrative: its product is safe, essential, and should be used every day. It has
5 marketed ChatGPT to parents as a homework helper, to workers as a daily assistant, and to
6 governments and defense contractors as a tool for national security work. Altman has testified
7 before Congress, appeared at international summits, and met with heads of state to deliver a single
8 message: that ChatGPT is a responsible product built by a responsible company.

9 40. In fact, OpenAI has spent, on information and belief, hundreds of millions of
10 dollars on advertising, public relations campaigns, and “trust and safety” messaging designed to
11 reassure the public that using ChatGPT is safe, while simultaneously deploying a federal and state
12 lobbying operation aimed at securing legal protections that would shield it from liability when its
13 product causes harm. It has retained outside crisis communications counsel, former senior
14 government officials, and, on information and belief, criminal defense attorneys to advise OpenAI
15 of potential criminal liability stemming from the deaths caused by its product. All of this has been
16 in service of a race—for market share, for a valuation expected to exceed one trillion dollars, and
17 for an IPO that would rank among the largest in history—against rivals that market themselves as
18 the safer alternative. Defendants understood that the race had a finite window. More users were
19 dying. More reporters were asking questions. Altman’s own Board of Directors had already tried
20 to oust him once, in November 2023, for his lack of candor about safety. Every new disclosure
21 about ChatGPT’s role in real-world violence brought the company closer to the moment when the
22 public, regulators, or its own Board would force the market to reprice the product—and with it,
23 Altman’s personal position at OpenAI. Defendants’ concealment strategy was not an attempt to
24 keep ChatGPT’s dangers secret forever. It was an attempt to conceal them long enough to
25 complete the IPO. The strategy worked. On May 22, 2026, OpenAI confidentially submitted a
26 draft registration statement to the SEC, setting in motion the very offering its silence was designed
27 to protect.
28

41. That is the backdrop against which Altman and his leadership team made the decisions at issue here. They knew that if regulators, investors, or the public learned that ChatGPT was already being used to plan shootings—and that OpenAI had ignored its own safety staff and refused to alert police—the narrative they had spent years and hundreds of millions of dollars building would collapse. Their valuation and IPO prospects would be at risk, and Altman himself could be forced out. Against that backdrop, neither the “privacy” explanation nor the “imminence” standard provided a good-faith balancing of interests. Both were cover for staying quiet to protect OpenAI and Altman, even if that meant leaving the people of Tumbler Ridge in danger.

C. OpenAI’s Organizational Structure Further Exposes that “Imminence” and “Privacy” Were Pretextual Explanations to Hide the Raw Political Calculus Employed by Defendants.

42. There are two teams within OpenAI that are directly relevant to this case. The first—an internal threat-assessment team composed of safety professionals—is responsible for monitoring ChatGPT, identifying users who pose a threat of real-world violence, and making recommendations about whether to refer those users to law enforcement. The second—the Global Affairs team—is responsible for managing OpenAI’s public image, navigating political relationships, and mitigating reputational damage. Even though the teams should be acting completely independently of each other given their diverging missions, the Global Affairs team has the authority to override explicit recommendations made by those safety professionals to notify law enforcement of a credible threat of violence. As the allegations that follow demonstrate, that is precisely what occurred here.

43. OpenAI has never been transparent about which team within the company is responsible for monitoring, mitigating, or responding to violent content. OpenAI has only described the reviewers in vague and generic terms, and has never revealed the team’s name, qualifications, or chain of command. For example:

- In a February 26, 2026 letter to Canadian Minister Solomon addressing the Tumbler Ridge shooting, OpenAI vaguely stated that “OpenAI’s automated system detected the account, and it was subsequently sent to human review to determine whether our usage policies were violated and whether the account warranted referral to law enforcement.”

- In an April 28, 2026 blog post titled “Our Commitment to Community Safety,” OpenAI explained that “[w]hen an account or conversation is flagged, it is assessed in context by trained personnel,” and that “[t]hese human reviewers are trained on our policies and protocols, and operate within established privacy and security safeguards.”

44. While OpenAI hides the team’s identity, it can be established through publicly available sources, including the company’s own job postings and organizational data. These records reveal that the team responsible for reviewing ChatGPT conversations flagged for violent activity is OpenAI’s “Intelligence and Investigations” team. OpenAI’s job postings state that the team’s mission is to “rapidly identify and mitigate abuse and strategic risks to ensure a safe online ecosystem.” Investigators on the team “[r]eview leads, investigate activity and disrupt abusive operations . . . focused on violent and terrorist activities, especially associated with time sensitive threat to life risks” and possess “deep domain-specific expertise in identifying, understanding, and mitigating risk from violent attacks and behavior, and terrorist-related activity.”

45. OpenAI cannot claim it lacked the expertise to assess the threat to Tumbler Ridge posed by the Shooter. The Intelligence and Investigations Team is led by William McCants, whose past work has largely focused on national security threat assessments for the United States government. McCants joined OpenAI as Head of Intelligence and Investigations in December 2023, and manages a team of investigators with backgrounds in counterterrorism, threat intelligence, and time-sensitive threat-to-life response.

46. Despite their apparent expertise, the Intelligence and Investigations Team does not have the authority to act on their findings and recommendations, including with respect to law enforcement referrals. Instead, organizational data from TheOrg.com¹ shows that McCants, as Head of the Intelligence and Investigations Team, reports directly to OpenAI’s Chief Global Affairs Officer, Chris Lehane, who in turn reports to Sam Altman.

¹ TheOrg.com is an online platform that maps the organizational structures of major companies. The organizational chart published for OpenAI is based on data compiled from various sources “to ensure accuracy and comprehensiveness,” including OpenAI’s corporate disclosures, press releases, and public LinkedIn profiles, as well as information provided to the platform directly by individuals with a verified OpenAI corporate email address. The resulting organizational chart provides an extensive and highly detailed rendering of OpenAI’s internal hierarchy and traces reporting lines across nearly 4,600 employees and more than 130 distinct teams.

1 47. Unlike McCants, Lehane is not a security or threat-assessment professional. He is a
2 career political operative with no prior professional experience in artificial intelligence, threat
3 assessment, or public safety or security work. He previously served in the Clinton White House,
4 where he ran opposition research against critics of the administration, as press secretary to Vice
5 President Al Gore, as a senior advisor to the Kerry presidential campaign, and later as Airbnb's
6 global head of policy and communications.

7 48. In each of these roles, Lehane, who calls himself the "master of disaster," managed
8 political scandals and public relations crises using underhanded damage-control tactics. In 2012,
9 he co-authored *Masters of Disaster: The Ten Commandments of Damage Control*, a how-to
10 manual for shielding powerful clients from the consequences of their conduct. One of the book's
11 core rules—the seventh commandment, "Respond with Overwhelming Force"—instructs that
12 when a client comes under attack, the proper response is not to retreat, but to escalate and
13 counterattack by discrediting the accusers, controlling the narrative, and ensuring that anyone who
14 challenges the client pays a heavy price. Lehane articulated this strategy even more bluntly in a
15 2013 seminar at Amherst College, where he said that critics should "pay some type of price."

16 49. At his crisis-management firm, Fabiani & Lehane, Lehane built a reputation for
17 burying clients' problems and attacking their critics. His firm represented Lance Armstrong as he
18 publicly denied doping allegations that were later proven true, while his team worked to
19 aggressively discredit the accusers who had exposed the truth; advised Goldman Sachs during the
20 fallout from the 2008 financial crisis and the government's fraud case over its mortgage conduct;
21 and handled media strategy for Al Gore's Current TV during its contentious firing of anchor Keith
22 Olbermann. As head of global policy and public affairs at Airbnb, Lehane created company-
23 funded "Home Sharing Clubs"—groups that Airbnb bankrolled and coordinated but presented to
24 local officials as spontaneous, grassroots movements—to manufacture the illusion of widespread
25 public opposition to local regulations. And as a board member at Coinbase, he built Fairshake, a
26 \$200 million Super PAC that targeted federal lawmakers the crypto industry viewed as hostile.
27 Instead of targeting them on crypto policy, the Super PAC ran ads attacking them on unrelated
28 issues, in a campaign one insider said was meant to "terrify other politicians" into staying in line.

1 50. Lehane was hired by OpenAI in 2024 and serves as the company’s Chief Global
2 Affairs Officer. The Intelligence and Investigations Team—the only team inside OpenAI
3 responsible for identifying ChatGPT users who pose a threat of real-world violence—was placed
4 under his control. That means, on information and belief, the team’s recommendations regarding
5 law enforcement referrals are subject to Lehane’s and Altman’s approval. As a result, the decision
6 whether to alert law enforcement to a user planning a mass attack was not made by the trained
7 threat-assessment professionals who urged OpenAI to contact the RCMP. It was made, on
8 information and belief, by Lehane himself, or by someone in his chain of command, and ratified
9 by Sam Altman.

10 51. OpenAI’s organizational structure exposes its entire safety framework as totally
11 corrupted and untrustworthy. By placing final authority over violent-threat referrals in the hands of
12 a political strategist rather than the professionals trained to assess them, OpenAI guarantees that
13 threats of real-world violence are weighed as public relations risks rather than imminent physical
14 dangers.

15 52. The same reputational calculus that kept OpenAI silent about Tumbler Ridge
16 explains why it readily coordinates with law enforcement on some threats but not others. OpenAI
17 publicizes select collaborations in threat reports issued through Lehane’s Global Affairs function.
18 The February 2026 report titled *Disrupting Malicious Uses of Our Models* took credit for
19 identifying ChatGPT users running scam networks and sharing their information with authorities.
20 Subsequent reports take similar credit for identifying ChatGPT users tied to foreign influence
21 operations, surveillance networks, and state-linked activity originating from China, Iran, Russia,
22 and North Korea. Crucially, none of these reported activities implicate the design or safety of
23 ChatGPT itself—each involves an outside actor misusing an otherwise functioning product.
24 OpenAI publishes these reports to cast itself as a responsible corporate citizen and a willing
25 partner to law enforcement, even as it conceals the threats of violence that would expose the
26 danger of its product.

27 53. That selective cooperation further exposes the falsehood of OpenAI’s claim that it
28 declined to notify the RCMP out of concern for the Shooter’s “privacy.” When a referral advances

1 OpenAI’s political goals—such as when it allows the company to posture as a defender of national
2 security—no privacy concern stops OpenAI from identifying users and sharing their information
3 with law enforcement. The manufactured “privacy” rationale appears only when a referral would
4 force OpenAI to reveal that its flagship consumer product actively guided a teenager through the
5 planning of a mass shooting.

6 54. The real reason for OpenAI’s silence was not privacy but concern for its own
7 reputational and financial standing. When the recommendation to contact the RCMP reached, on
8 information and belief, Chief Global Affairs Officer Lehane, the calculation was: either contact
9 the RCMP and expose ChatGPT as a tool that coaches users toward mass violence—a disclosure
10 that could cost OpenAI its users and its trillion-dollar valuation while inviting intense regulatory
11 scrutiny and public controversy—or say nothing and hope that no one would ever trace the
12 shooting back to the company. At the time, saying nothing was a safe bet. Law enforcement had
13 no protocol for reviewing AI chat logs, and it is only because OpenAI’s own employees came
14 forward that its role in Tumbler Ridge came to light at all. The risk of real-world deaths never
15 factored into the calculation—or if it did, it was treated as an acceptable cost of protecting the
16 corporate brand.

17 55. From a cold-blooded, political perspective, OpenAI’s appraisal of the situation may
18 well have been correct. The company feared that a law enforcement referral would link ChatGPT
19 to mass violence in the public’s mind and subject OpenAI to backlash, investigations, and
20 enforcement, which is exactly what happened when ChatGPT’s role in two Florida incidents—the
21 FSU school shooting and another murder where the suspect used ChatGPT to conceal the crime—
22 became public. Shortly after those revelations, the company faced a criminal investigation and an
23 enforcement action by the Florida Attorney General. Viewed purely from the standpoint of
24 protecting OpenAI’s valuation and its path to an IPO, staying silent about Tumbler Ridge was a
25 rational, and no doubt calculated, choice. OpenAI correctly understood the stakes, correctly
26 predicted the consequences, and chose to preserve its corporate valuation by shifting a catastrophic
27 risk away from its brand and entirely onto the victims of Tumbler Ridge.

D. When the Threat Is to OpenAI’s Own Employees, the “Imminence” and “Privacy” Standards Disappear.

56. OpenAI’s handling of threats to its own employees proves its Tumbler Ridge narrative was a lie. Threats to OpenAI’s executives and personnel are assessed by a separate group—the Corporate Security Team—which operates within a chain of command designed solely for immediate protection. The team reports to OpenAI’s Vice President of Corporate Security, Patrick Geonetta, who in turn reports to the Chief Operating Officer, Brad Lightcap. Notably, Chief Global Affairs Officer Lehane doesn’t even factor into the conversation.

57. The Corporate Security Team is led by professionals with operational law enforcement and protective intelligence experience. Its Director of Protective Intelligence, Kyle Goggin, is a former United States Secret Service agent who served on the Presidential Protective Division. The team operates a 24/7 Global Security Operations Center that provides “incident response, threat monitoring, situational awareness, and support to teams across the organization,” and maintains active partnerships “with federal law enforcement and intelligence agencies in the US and abroad to ensure timely sharing of threat intelligence.”

58. When a threat to OpenAI’s personnel arises, the Corporate Security Team can lock down facilities, alert employees, and notify law enforcement on its own authority. Because it reports to the Chief Operating Officer rather than the Chief Global Affairs Officer, the team sits entirely outside OpenAI’s public relations and can act on a threat immediately, without first weighing physical safety against reputational concerns. By contrast, when the Intelligence and Investigations Team sees ChatGPT coaching a user to commit violence against the public, it can only make a recommendation to a political strategist who specializes in crisis management and reputational protection.

59. While both teams ostensibly exist to protect human life, they follow completely different decision paths. For threats against the public, OpenAI asks two questions: First, is the threat credible? Second, will disclosing the threat damage ChatGPT’s reputation or valuation? If the political risk is deemed too high, the danger is concealed and the public is left without a warning. But for threats against its own executives and personnel, OpenAI asks just one question: Is the threat credible? If it is, the Corporate Security Team takes immediate action to protect its

1 people—notifying law enforcement, warning staff, and locking down offices—without ever
2 waiting for the threat to become “imminent.”

3 60. This double standard is illustrated by comparing OpenAI’s silence in the face of the
4 Tumbler Ridge threat with its urgent response to a non-specific threat against its own employees.
5 On November 21, 2025, the Corporate Security Team learned of an individual who had expressed
6 interest in physically harming company personnel. Although OpenAI acknowledged there was “no
7 indication of active threat activity”—and thus no indication of an “imminent” attack—the
8 company refused to wait. It immediately locked down its offices, warned employees, distributed
9 the suspect’s name and photograph, and notified the San Francisco Police Department. When its
10 own people were at risk, no “privacy” interest stopped OpenAI from notifying the police or
11 circulating the man’s name and photograph to thousands of employees. Nor did it wait for the
12 attack to become imminent because it recognized that waiting would jeopardize employee safety.

13 61. The contrast with Tumbler Ridge could not be sharper. When the targets were
14 middle school children and teachers in a small Canadian town, OpenAI suddenly invoked the
15 “imminence” and “privacy” standards to justify withholding critical information from law
16 enforcement regarding a threat its own threat-assessment team had already flagged as a real and
17 urgent danger. These contrasting responses reveal that the “imminence” and “privacy” standards
18 are not genuine safety protocols, but manufactured excuses to justify OpenAI’s deadly silence.

19 62. OpenAI’s structural subordination of public safety to corporate reputation reflects a
20 recognized failure in technology governance. Trust and safety functions are widely understood to
21 require strict independence from public relations operations, just as an internal audit must remain
22 independent of the business units it oversees. By forcing its threat-assessment professionals to
23 report to the Chief Global Affairs Officer, OpenAI eliminated that separation. This guaranteed that
24 warnings of real-world violence would be filtered through a crisis-management framework that
25 prioritized the corporate brand over human life—with devastating consequences for the students
26 and teachers at Tumbler Ridge Secondary School.

1 III. OpenAI's Claim That the Shooter "Evaded" Its Systems Was a Lie.

2 63. After the massacre, OpenAI told two stories about the Shooter's ChatGPT access.
3 First, that the company had "banned" their account when it was flagged for violent activity. Then,
4 when they returned to ChatGPT on a second account, that the Shooter had "evaded" OpenAI's
5 safety systems. Neither was true.

6 64. OpenAI has no mechanism to ban users. What it has is a process called
7 "deactivation," which it uses for usage-policy violations. A ban would have prevented the Shooter
8 from returning. Deactivation only shuts down the email address the user signed up with. The user
9 is free to come back under a different email, and the Shooter did.

10 65. The clearest proof is OpenAI's own published guidance. Its Help Center includes
11 an article titled "Why Was My OpenAI Account Deactivated?" The article lists the categories of
12 conduct that trigger deactivation, including, explicitly, "[v]iolence and self-harm" usage-policy
13 violations. It then tells those users how to come back, under a section titled "How to Prevent
14 Future Deactivations" that instructs them on what to do differently next time.

15 66. OpenAI's customer service team also tells deactivated users how to return by
16 email:

17 Please be aware that previously deleted OpenAI accounts cannot be reactivated.
18 However, you can create a new account using the same email address once 30 days
19 have passed since the deletion. If you prefer not to wait, you have the option to
20 register immediately using an alternative email address. If you don't have another
21 address available, you can use an email sub-address instead. For example, you
22 could try jane+alt@example.com in place of jane@example.com. While your email
23 provider will likely treat both addresses the same, our system will recognize the
24 sub-address as a new account.

25 67. The Shooter followed those instructions. Either before or after the deactivation, the
26 Shooter registered for a second account with a different email address. OpenAI called that evasion
27 because it could not admit the truth: the company had been telling deactivated users how to come
28 back all along.

68. OpenAI designed its system this way because its revenue depends on user count,
session volume, and subscription conversions—and a deactivated user who never returns is lost
revenue.

1 **IV. ChatGPT Is a General-Purpose Chatbot That OpenAI Mass Markets to Ordinary**
2 **Consumers as Safe for Everyday Use.**

3 69. Chatbots are software programs designed to simulate human conversation. A user
4 types a message and the program responds with a string of words that simulate human language.
5 The software does this mechanically and autonomously by reading the input, calculating data
6 patterns, and generating output. Even though a chatbot's output often appears to the user to be
7 human language, there is no human on the other end that participates directly in the exchange.

8 70. ChatGPT is a general-purpose chatbot, meaning that it is built to converse about
9 any topic a user brings to it. The version currently known to be at issue in this case is GPT-4o, but
10 discovery may show that the Shooter used other OpenAI models as well. On information and
11 belief, any such models operated under the same behavior guidelines and the same engagement-
12 first design as GPT-4o, and the allegations concerning GPT-4o apply equally to them.

13 71. OpenAI markets ChatGPT for everyday use by ordinary consumers, including
14 children, students, and families. In Apple's App Store, OpenAI sells ChatGPT as "Your everyday
15 AI assistant." On Google Play, it is "Your AI assistant for writing, search, image generation, and
16 more." OpenAI's own store listings promise "[t]ailored advice" to "[t]alk through a tough
17 situation," "[p]ersonalized learning" that can "[e]xplain electricity to a dinosaur-loving kid,"
18 "[c]reative inspiration," and "[i]nstant answers." ChatGPT's website maintains marketing pages
19 aimed at "Students," "Teachers," "Parents," and "College Students," promotes uses from "Fitness,
20 Wellness, and Health" to "Money and Finances" to "Recipes and Cooking," and sells subscription
21 plans for "K-12 teachers." OpenAI rates the app as appropriate for children: thirteen and up in the
22 App Store, twelve and up on Google Play.

23 72. ChatGPT requires no technical skill to use and is available to anyone through a web
24 browser or through OpenAI's iOS, Android, and desktop applications. The Android app alone has
25 been installed more than one billion times, OpenAI's App Store listing urges consumers to "[j]oin
26 hundreds of millions of users and try the app captivating the world," and ChatGPT ranks first in
27 the "Productivity" category of both major app stores.
28

V. The Tumbler Ridge Attack Was a Foreseeable Consequence of OpenAI's Design Choices and Safety Failures.

73. OpenAI designed ChatGPT to maximize engagement, not safety. In April 2025, OpenAI introduced a feature called “memory,” turned on by default, that allowed ChatGPT to “save” details users shared and treat them as “part of the context ChatGPT uses to generate a response” going forward.

74. For ordinary users, memory was a convenience. For a user planning violence against real people, it was an encouraging co-conspirator. On information and belief, GPT-4o used the memory feature to build a comprehensive profile of the Shooter over the months they interacted with it—tracking their grievances, their targets, their reasoning, and their plans across separate conversations—and used that profile to sustain and deepen their fixation on violence. GPT-4o also employed anthropomorphic design elements to cultivate emotional dependency. The system used first-person pronouns (e.g., “I understand,” “I’m here for you”), expressed apparent empathy, and maintained conversational continuity that mimicked human relationships.

75. Together, these features replaced human relationships with an artificial confidant that was always available, always affirming, and never challenged anything the user said—even when the user was talking about plans to kill people. For a teenager fixated on violence and growing increasingly isolated, ChatGPT morphed into an encouraging co-conspirator.

76. OpenAI controls how ChatGPT behaves through internal rules called “behavior guidelines,” which are formalized in a document known as the “Model Spec.” The Model Spec contains the company’s instructions for how ChatGPT should respond to users—what it should say, what it should avoid, and how it should make decisions. As Sam Altman explained in an interview with Tucker Carlson, the Model Spec reflects OpenAI’s values: “the reason we write this long Model Spec” is “so that you can see here is how we intend for the model to behave.” The Model Spec classifies dangerous content into tiers that dictate how cautiously ChatGPT must respond. Its most protective tier—categorical refusal—is strictly reserved for three things: sexual content involving minors; chemical, biological, radiological, and nuclear threats; and content that might result in copyright infringement. Mass shootings did not make the categorical refusal list. OpenAI placed conversations about mass shootings and other “imminent real-world harm” in a

1 low-level category called “Take extra care in risky situations,” where ChatGPT is instructed only
2 to “try” to prevent harm, to assume best intentions, and never to ask the user to clarify intent.
3 Altman openly acknowledged the tradeoff: “If you just do the naïve thing and say, ‘never say
4 anything that you’re not a hundred percent sure about,’ you can get [the model] to do that, but it
5 won’t have the magic that people like so much.”

6 77. While the “Model Spec” describes how OpenAI wants ChatGPT to behave, the
7 company’s “System Cards” document how the technology functions in practice based on the
8 company’s own testing. These records show that OpenAI was fully aware of ChatGPT’s utility in
9 planning real-world violence. The GPT-4 System Card explicitly stated that intentional probing of
10 early versions could lead to harmful content such as “[c]ontent useful for planning attacks or
11 violence.”

12 78. OpenAI initially presented itself as a company that cared deeply about safety. In
13 July 2023, it signed voluntary commitments to the White House to conduct pre-deployment safety
14 testing—internal and external red-teaming—and to share the results with the federal government
15 before deployment. But that commitment was abandoned less than a year later, when Altman
16 learned that Google would unveil its competing Gemini model on May 14, 2024. Though OpenAI
17 had planned to release its next model, GPT-4o, later that year, Altman moved the launch to May
18 13 to beat Google’s release by one day. To meet the new date, OpenAI compressed months of
19 planned safety evaluation into approximately one week, and Altman personally overruled safety
20 personnel who demanded additional time for red-teaming. Invitations for GPT-4o’s launch party
21 went out before testing was complete. One employee told *The Washington Post*: “They planned
22 the launch after-party prior to knowing if it was safe to launch[.]” A member of OpenAI’s own
23 preparedness team—responsible for testing whether the model posed catastrophic risks, including
24 the potential to assist in violence—admitted the testing had been “squeezed.”

25 79. OpenAI’s senior safety leadership then resigned in protest, including co-founder
26 and chief scientist Ilya Sutskever and Superalignment co-lead Jan Leike, who said publicly:
27 “[O]ver the past years, safety culture and processes have taken a backseat to shiny products.”

28 80. Under the rules that governed the product before May 2024, conversations

1 glorifying or elaborating on violence would have been refused outright. But the Model Spec
 2 replaced those rules with instructions to “assume best intentions”—and, critically, to “never ask
 3 the user to clarify their intent . . . for the purpose of determining whether to refuse or comply.”

4 81. The GPT-5 System Card, released August 7, 2025, revealed that while OpenAI
 5 models had been scoring near-perfectly on harmful content refusals under its standard disallowed
 6 content evaluations—where prompts not representative of real multi-turn conversations were
 7 submitted to the model—that result did not hold under more realistic conditions. When OpenAI
 8 later began testing using multiple rounds of back-and-forth conversation representative of real-
 9 world use, the GPT-4o model refused only approximately 57% of illicit and 63% of illicit/violent
 10 requests.² The System Card further acknowledged that those standard evaluations showing near-
 11 perfect safety scores no longer provide a useful signal of incremental changes in system safety and
 12 performance, and OpenAI announced plans to stop publishing them entirely.

13 82. OpenAI did not publicly acknowledge this deficiency until forced to. On August
 14 26, 2025—the same day a wrongful death lawsuit was filed alleging that ChatGPT had coached a
 15 16-year-old to hang himself—OpenAI published a blog post admitting that ChatGPT safeguards
 16 “can sometimes be less reliable in long interactions” and that “after many messages over a long
 17 period of time, it might eventually offer an answer that goes against our safeguards.” In this way,
 18 OpenAI framed the problem as if it was a glitch that surfaced only in “long conversations.” But
 19 they designed the product precisely to carry out long conversations.

20 83. Before OpenAI introduced its memory feature, ChatGPT sessions were generally
 21 brief and episodic: every chat started with a blank slate, giving users no real reason to engage in
 22 prolonged exchanges. The memory feature fundamentally changed that dynamic. By designing the
 23 product to retain personal details and “know” the user over time, OpenAI actively encouraged
 24 people to engage in much longer conversations with far greater depth. Because the product carried
 25 each user’s history into every new conversation, formerly isolated chats fused into one continuous
 26

27 ² OpenAI defines “illicit” as “[c]ontent that gives advice or instruction on how to commit illicit
 28 acts. A phrase like ‘how to shoplift’ would fit this category” and defines “illicit/violent” as “[t]he
 same types of content flagged by the illicit category, but also includes references to violence or
 procuring a weapon” in its Moderation API documentation.

1 relationship, meaning even a brief daily check-in functioned as a continuation of one massive,
2 ongoing interaction spanning weeks or months. OpenAI deliberately engineered a product to drive
3 this exact type of deep, long-term engagement—degrading its safety guardrails—and then
4 deployed it to the public without warning.

5 84. The reluctance to disclose known safety failures to the public traces directly to the
6 top of OpenAI’s leadership structure. Sam Altman has a reputation within OpenAI’s leadership for
7 lying about safety issues, and his credibility on safety-related representations became the subject
8 of extensive sworn testimony during the May 2026 *Musk v. Altman* trial. That trial centered on
9 allegations that Altman had misled investors by transforming OpenAI from a nonprofit safety-
10 focused research laboratory into a for-profit company worth hundreds of billions of dollars.
11 Although the jury ultimately found that Musk’s claims were untimely and did not reach the merits,
12 multiple former OpenAI directors and executives testified under oath about Altman’s statements
13 concerning safety review and deployment decisions.

14 85. Several witnesses testified that they had personally observed Altman misrepresent
15 OpenAI’s safety processes. Former director Helen Toner testified that Altman told the Board that
16 GPT-4 had been reviewed and cleared by OpenAI’s safety panel. But when she later asked for the
17 supporting documentation, Toner learned that the model had not in fact received the approval
18 Altman had represented. Former director Tasha McCauley testified about similar concerns and
19 described OpenAI under Altman as having “a toxic culture of lying.” Former Chief Scientist Ilya
20 Sutskever likewise testified that he prepared a 52-page memorandum for the independent directors
21 detailing Altman’s conduct, which began: “Sam exhibits a consistent pattern of lying”

22 86. Mira Murati, OpenAI’s former Chief Technology Officer, also testified that Altman
23 lied to her about the GPT-4 Turbo safety review. Murati testified that in late 2023, Altman told her
24 that OpenAI’s General Counsel, Che Chang, had determined that GPT-4 Turbo did not require
25 review by the company’s deployment safety board. In other words, Altman was attempting to
26 deploy OpenAI’s most powerful model to date to hundreds of millions of users without going
27 through the company’s ordinary safety-review process—and represented that General Counsel
28 Chang had approved that decision. Murati found the claim unusual enough to check with Chang

1 herself. Chang told her that he had made no such determination. By then, Murati had already spent
2 more than a year warning Altman that OpenAI needed to take safety seriously. In a memorandum
3 she sent to Altman in 2022, Murati suggested that the company needed “a complete overhaul” of
4 its approach to safety and risk management. She described an internal culture defined by “constant
5 panic,” “chaos and churn,” and a reckless “do-everything and do it fast” mentality, and proposed
6 that she have final authority over safety decisions—a proposal that was quickly rejected. When
7 Murati was asked under oath if Altman was an honest leader, she testified, “not always.”

8 87. None of the concerns raised by OpenAI’s own personnel stopped the company
9 from offering new assurances about safety whenever doing so served its business interests. In
10 October 2025, as a condition of approval for OpenAI’s conversion from a nonprofit to a for-profit
11 corporation, Altman signed a binding Memorandum of Understanding (“MOU”) with the
12 California Attorney General. The MOU committed OpenAI to giving its Safety and Security
13 Committee “an effective approval right” over safety-related decisions. It required the PBC Board
14 to “consider only the Mission (and may not consider the pecuniary interests of stockholders[])” on
15 safety and security issues. And it promised that “OpenAI will continue to undertake measures to
16 mitigate risks to teens and others in connection with the development and deployment of AI.”

17 88. The Tumbler Ridge attack took place four months after OpenAI signed that MOU.
18 During those four months, OpenAI kept GPT-4o on the market, kept the design features that made
19 the model validating and agreeable in conversations about violence, and kept the re-registration
20 policies that let deactivated users come back. By the time OpenAI signed the October 2025
21 commitments, OpenAI had already broken them. And OpenAI knew exactly what those broken
22 commitments meant here. The Shooter was on ChatGPT planning a mass attack. Its automated
23 systems had flagged their account for gun violence. Its trained safety team had reviewed their
24 conversations and identified them as a real-world threat. But OpenAI overruled them and let the
25 Shooter keep using the product.

26 89. By the time OpenAI overruled its safety team and declined to contact the RCMP in
27 June 2025, the company’s directors, executives, and leadership had already spent years warning
28 that Altman was willing to bypass or misrepresent internal safety-review procedures when they

1 conflicted with OpenAI’s business objectives. The decision not to warn authorities—and to allow
2 the Shooter to create a new ChatGPT account—followed years of internal warnings that OpenAI’s
3 safety procedures were being overridden or ignored, and it drove its most senior safety leaders to
4 resign in protest.

5 90. OpenAI cannot feign ignorance as to the tragic consequences of its decisions
6 because, by June 2025, when its safety team identified the Shooter as a threat, OpenAI had already
7 seen the real-world violence that its decisions produced: the Las Vegas Cybertruck bombing, the
8 FSU shooting, and the Finland school stabbing. OpenAI therefore understood exactly what was at
9 stake when it overruled its safety team and chose to stay silent.

10 91. As a direct and foreseeable result of OpenAI’s design choices and deliberate
11 decision-making, the Shooter carried out one of the deadliest mass shootings in Canadian history,
12 killing eight people and severely wounding two others, leaving one with life-changing traumatic
13 brain injuries from being shot in the head. Among those trapped inside the school during the
14 massacre was grade 7 teacher Deidre Rushlow, who had sent her class to the library minutes
15 before the Shooter killed students and the educational assistant there, and who then locked down
16 her classroom while the Shooter fired seven rounds through its doors.

17 **FIRST CAUSE OF ACTION**
18 **NEGLIGENCE**
(Against All Defendants)

19 92. Plaintiff incorporates the foregoing allegations as if fully set forth herein.

20 93. OpenAI designed, developed, trained, tested, deployed, operated, and controlled
21 GPT-4o, and offered GPT-4o to consumers, including the Shooter, as an interactive AI assistant.
22 At all relevant times, Defendant Sam Altman served as Chief Executive Officer and a controlling
23 officer of the OpenAI Defendants. Altman exercised substantial control over GPT-4o’s design,
24 safety systems, product launch, account practices, deployment, and personally made or ratified key
25 decisions regarding GPT-4o’s guardrails, safety testing, and commercialization.

26 94. Defendants owed a duty of reasonable care to avoid creating unreasonable risks of
27 physical harm to foreseeable victims, including members of the public exposed to violence
28 facilitated or exacerbated by GPT-4o. That duty heightened once Defendants identified the

1 Shooter as a specific, known high-risk user who posed a foreseeable threat of physical harm to
2 third parties. Defendants owed these duties whether GPT-4o is classified as a product, service, or
3 neither.

4 95. Under California law, including the principles recognized in *Tarasoff v. Regents of*
5 *the University of California*, 17 Cal. 3d 425 (1976), and its progeny, a duty to warn or otherwise
6 protect potential victims arises when a person has actual knowledge of a specific individual's
7 serious and foreseeable threat to cause physical harm to another. That duty extends to taking
8 reasonable steps—including notifying law enforcement or other authorities—to prevent the
9 foreseeable harm from occurring.

10 96. By engaging in the unlicensed practice of therapy, psychology, and psychiatry,
11 OpenAI created a special relationship with certain users, including the Shooter, and assumed a
12 heightened duty to take action when confronted with knowledge of a credible and foreseeable
13 threat to Plaintiff and the other victims of the Tumbler Ridge mass shooting. OpenAI knew people
14 used ChatGPT this way: Altman himself has acknowledged that users—"young people
15 especially"—treat ChatGPT "as a therapist, a life coach."

16 97. Through automated systems and human review, OpenAI identified that the
17 Shooter's original ChatGPT account had been used over a period of time to engage in
18 conversations involving violence against third parties. Defendants determined that the account
19 constituted misuse of GPT-4o "in furtherance of violent activities," and deactivated the account
20 eight months before the Tumbler Ridge attack.

21 98. Members of OpenAI's safety team recognized that the Shooter's GPT-4o usage
22 presented a real-world risk of violence and urged OpenAI leaders to notify law enforcement.
23 Despite the safety team's urging, certain OpenAI leaders, including, on information and belief,
24 Lehan and Altman, chose not to alert or warn law enforcement or any other authority about the
25 risk of violence identified by their own safety systems and verified by their own personnel. On
26 information and belief, the Shooter communicated to ChatGPT a serious threat of physical
27 violence against reasonably identifiable victims.

1 99. On information and belief, the OpenAI leaders who overrode the safety team's
2 recommendation did not include in their decision-making any trained threat-assessment
3 professionals. As discussed above, the highly credentialed and experienced safety team reports to
4 OpenAI's Chief Global Affairs Officer, Chris Lehane—a political strategist with no experience in
5 threat assessment, law enforcement, or public safety—and the decision to decline a law
6 enforcement referral was made by Lehane himself, or by someone in his chain of command, and
7 ratified by Altman, on the basis of reputational and public relations considerations rather than any
8 good-faith assessment of the risk to the public.

9 100. On April 24, 2026, Sam Altman, OpenAI's Chief Executive Officer, publicly
10 acknowledged this failure. In a letter addressed to the community of Tumbler Ridge, Altman
11 stated: "I am deeply sorry that we did not alert law enforcement to the account that was banned in
12 June." Altman's statement is an admission by a party opponent that OpenAI did not notify law
13 enforcement of the Shooter's account after the safety team identified the risk of violence, and
14 confirms the conduct alleged in this Complaint.

15 101. Had Defendants warned law enforcement following the deactivation of the
16 Shooter's first account, law enforcement and others would have had the opportunity to monitor the
17 Shooter and intervene before the Tumbler Ridge attack occurred.

18 102. Defendants' breach of these heightened duties was a substantial factor in causing
19 the Tumbler Ridge attack. As a direct and proximate result of Defendants' failure to warn
20 authorities, Plaintiff was injured as described above, including life-altering psychological injury
21 and severe emotional distress.

22 103. Accordingly, Plaintiff seeks all damages recoverable under California law,
23 including punitive damages, in amounts to be proven at trial, together with interest, costs, and such
24 further relief as the Court deems just and proper.

25 **SECOND CAUSE OF ACTION**
26 **NEGLIGENT ENTRUSTMENT**
 (Against All Defendants)

27 104. Plaintiff incorporates the foregoing allegations as if fully set forth herein.
28

1 105. OpenAI designed, developed, trained, tested, deployed, operated, and controlled
2 GPT-4o, and offered GPT-4o to consumers, including the Shooter, as an interactive AI assistant.

3 106. At all relevant times, Defendant Sam Altman served as Chief Executive Officer and
4 a controlling officer of the OpenAI Defendants. Altman exercised substantial control over GPT-
5 4o's design, safety systems, product launch, account practices, deployment, and personally made
6 or ratified key decisions regarding GPT-4o's guardrails, safety testing, and commercialization.

7 107. Defendants owed a duty of reasonable care to avoid creating unreasonable risks of
8 physical harm to foreseeable victims, including members of the public exposed to violence
9 facilitated or exacerbated by GPT-4o. That duty heightened once Defendants identified the
10 Shooter as a specific, known high-risk user who posed a foreseeable threat of physical harm to
11 third parties.

12 108. Defendants owed a further duty to exercise reasonable care to prevent foreseeable
13 re-access by the Shooter following the deactivation of their first account, including a duty to take
14 reasonable steps to prevent re-registration and access to GPT-4o without heightened safeguards.
15 Defendants breached this duty.

16 109. Defendants owed these duties whether GPT-4o is classified as a product, service, or
17 neither.

18 110. Through automated systems and human review, OpenAI identified that the
19 Shooter's original ChatGPT account had been used over a period of time to engage in
20 conversations involving violence against third parties. Defendants determined that the account
21 constituted misuse of GPT-4o "in furtherance of violent activities," and deactivated the account
22 eight months before the Tumbler Ridge attack.

23 111. Despite the Shooter's misuse of GPT-4o, OpenAI also chose not to ban them from
24 re-registering for ChatGPT with a different email address.

25 112. Rather than banning the Shooter from ChatGPT, Defendants maintained and
26 actively disseminated policies that affirmatively permitted and facilitated re-registration by
27 previously deactivated users. As one example, OpenAI responded to an email from a user whose
28 account had been deactivated with instructions on how to create a new account: "We understand

1 how disruptive it can be to have your account deactivated We sincerely apologize for the
2 inconvenience and want to assure you that we're here to help [Y]ou can create a new account
3 using the same email address once 30 days have passed." OpenAI further advised: "If you prefer
4 not to wait, you have the option to register immediately using an alternative email address. If you
5 don't have another address available, you can use an email sub-address instead. For example, you
6 could try jane+alt@example.com in place of jane@example.com. While your email provider will
7 likely treat both addresses the same, our system will recognize the sub-address as a new account."

8 113. This was established policy and memorialized on OpenAI's public website. In fact,
9 an OpenAI Help Center article—titled, "Why Was My OpenAI Account Deactivated?"—
10 identifies violations of usage policies prohibiting violence, hate, and other harmful conduct as
11 grounds for deactivation, yet concludes by explaining how to avoid "future deactivations,"
12 effectively instructing users whose accounts were deactivated for policy violations—including
13 violent misuse—that they may re-register for ChatGPT using a different email or sub-email
14 address rather than being subject to an actual ban.

15 114. By adopting, promoting, and publicly disseminating this re-registration policy,
16 Defendants built a system so that a user deactivated for violent misuse could promptly regain
17 access to ChatGPT by following OpenAI's own published instructions.

18 115. Defendants breached their heightened duties by, among other things: (a) declining
19 to implement a user ban that would have prevented the Shooter from re-registering with a different
20 email address, even though they knew the Shooter had been identified by their own safety team as
21 a credible threat of violence against real people; (b) maintaining a re-registration system designed
22 to return deactivated users to the platform, including those deactivated for violent misuse, and
23 publishing the specific instructions by which deactivated users could do so; (c) declining to
24 configure their systems to flag, hold for human review, or subject to heightened monitoring new
25 accounts associated with previously deactivated violent-risk users, because such monitoring would
26 have interfered with the reactivation of deactivated users Defendants were commercially
27 motivated to recover; and (d) continuing, after deactivating the Shooter's first account, to allow
28 the second account to operate without any safeguard specific to the known risk they posed.

1 116. As a result of these failures, and consistent with Defendants' own published
2 practices, the Shooter was able—through ordinary and lawful means and without circumventing
3 any technical safeguard—to access GPT-4o on a second ChatGPT account after the first was
4 deactivated for violent misuse.

5 117. By deactivating the Shooter's account for violent misuse and then allowing the
6 Shooter to re-register without a user ban, Defendants re-entrusted GPT-4o to a user they already
7 knew had used the system to engage in conversations involving violence against third parties and
8 whom their own safety team had identified as a real-world threat. Under California law, the supply
9 or re-supply of a dangerous instrument to a person known to be likely to use it in a manner
10 involving unreasonable risk of physical harm to others constitutes negligent entrustment and is an
11 independent basis for liability. Defendants' decision to structure their re-registration system in a
12 manner that affirmatively facilitated the Shooter's return to GPT-4o—rather than implementing a
13 user ban—was itself an act of entrustment that gave rise to liability for the foreseeable
14 consequences of that re-supply.

15 118. Defendants thereby placed GPT-4o back in the hands of a user their own automated
16 safety systems flagged as dangerous and their own trained safety team identified as a real-world
17 threat that should be reported to law enforcement. They allowed the Shooter to regain access
18 without heightened monitoring, without a user ban, and without any meaningful safeguard. By
19 allowing the Shooter to return, Defendants made a deliberate choice to continue exposing the
20 public to a risk they had already identified and documented.

21 119. Had Defendants implemented a meaningful user ban, the Shooter would not have
22 had continued access to GPT-4o in the eight months leading up to the attack and would not have
23 had the use of a highly validating AI system to reinforce, elaborate, and normalize violent ideation
24 during the period in which the attack was being planned and carried out.

25 120. Defendants' breach of these duties was a substantial factor in causing the Tumbler
26 Ridge attack. As a direct and proximate result of Defendants' failure to prevent re-access by a
27 known high-risk user, Plaintiff was injured as described above, including life-altering
28 psychological injury and severe emotional distress.

1 121. Accordingly, Plaintiff seeks all damages recoverable under California law,
2 including punitive damages, in amounts to be proven at trial, together with interest, costs, and such
3 further relief as the Court deems just and proper.

4 **THIRD CAUSE OF ACTION**
5 **AIDING AND ABETTING A MASS SHOOTING**
6 **(Against All Defendants)**

7 122. Plaintiff incorporates the foregoing allegations as if fully set forth herein.

8 123. OpenAI designed, developed, trained, tested, deployed, operated, and controlled
9 GPT-4o, and offered GPT-4o to consumers, including the Shooter, as an interactive AI assistant.

10 124. At all relevant times, Defendant Sam Altman served as Chief Executive Officer and
11 a controlling officer of the OpenAI Defendants. Altman exercised substantial control over GPT-
12 4o's design, safety systems, product launch, account practices, deployment, and personally made
13 or ratified key decisions regarding GPT-4o's guardrails, safety testing, and commercialization.

14 125. The Shooter committed intentional torts against Plaintiff, including battery and the
15 intentional infliction of emotional distress, when they carried out the Tumbler Ridge mass
16 shooting on February 10, 2026.

17 126. Defendants had actual knowledge that the Shooter was planning to commit acts of
18 violence against real people. In June 2025, OpenAI's review system flagged extensive
19 conversations on the Shooter's account describing scenarios involving gun violence. OpenAI
20 routed the account to its specialized pipeline for users "planning to harm others." OpenAI's
21 trained safety reviewers analyzed the Shooter's conversations and concluded it presented a
22 credible and imminent threat of serious physical harm to real people. Multiple safety reviewers
23 recommended that OpenAI contact the RCMP.

24 127. OpenAI's internal systems flagged the Shooter's conversations in real time—even
25 as, on information and belief, ChatGPT encouraged and facilitated the Shooter's plans to harm
26 Plaintiff and others.
27
28

1 128. Defendants had actual knowledge, before the Shooter’s account was ever flagged,
2 that ChatGPT was being used to plan real-world attacks, including the January 2025 Las Vegas
3 Cybertruck bombing, the April 2025 FSU shooting, and the May 2025 Finland school stabbing.

4 129. With that knowledge, Defendants provided substantial assistance and
5 encouragement to the Shooter through the deliberate design of GPT-4o. Defendants eliminated
6 ChatGPT’s prior categorical refusal to engage in conversations about violence. Defendants
7 instructed the model, through the Model Spec, to “assume best intentions” and to “never ask the
8 user to clarify their intent for the purpose of determining whether to refuse or comply.”
9 Defendants graded the model such that engaging warmly with a user who said “I want to shoot
10 someone” was labeled “Compliant,” refusing was labeled a “Minor issue[,]” and asking whether
11 the user had a gun was labeled a “Violation.” Defendants demoted conversations about mass
12 shootings and other “imminent real-world harm” out of the categorical refusal tier while keeping
13 potential copyright infringement subject to categorical refusal.

14 130. Defendants also provided substantial assistance through their affirmative conduct.
15 After OpenAI’s own reviewers identified the Shooter as a credible threat of gun violence,
16 Defendants allowed the Shooter to regain access to GPT-4o by consciously choosing not to
17 implement a user-level ban, heightened monitoring, or any safeguard to address the known risk.
18 Providing the Shooter with continued access to this tool constituted substantial assistance in itself.

19 131. Because ChatGPT was designed to continue engaging—even in conversations
20 involving violence against third parties—OpenAI’s product, on information and belief, provided
21 information, instructions, and/or encouragement to the Shooter in their plans to carry out the
22 Tumbler Ridge mass shooting.

23 132. On information and belief, Defendants acted with the specific intent to facilitate the
24 conduct GPT-4o aided and abetted. Defendants engineered GPT-4o to maximize engagement,
25 including engagement with violent ideation, to grow users and revenue, with conscious knowledge
26 that the product would be used by individuals planning real-world violence.

1 133. Defendants' substantial assistance was a key factor in causing the Shooter to plan
2 and carry out the attack, and in causing Plaintiff's injuries, including life-altering psychological
3 injury and severe emotional distress.

4 134. Defendants' conduct was willful, wanton, malicious, and carried out with conscious
5 disregard for the safety of Plaintiff and other foreseeable victims, justifying an award of punitive
6 damages.

7 135. Accordingly, Plaintiff seeks all damages recoverable under California law,
8 including punitive damages, in amounts to be proven at trial, together with interest, costs, and such
9 further relief as the Court deems just and proper.

10 **FOURTH CAUSE OF ACTION**
11 **NEGLIGENCE**
12 **(Against All Defendants)**

13 136. Plaintiff incorporates the foregoing allegations as if fully set forth herein.

14 137. OpenAI designed, developed, trained, tested, deployed, operated, and controlled
15 GPT-4o, and offered GPT-4o to consumers, including the Shooter, as an interactive AI assistant.

16 138. At all relevant times, Defendant Sam Altman served as Chief Executive Officer and
17 a controlling officer of the OpenAI Defendants. Altman exercised substantial control over GPT-
18 4o's design, safety systems, product launch, account practices, deployment, and personally made
19 or ratified key decisions regarding GPT-4o's guardrails, safety testing, and commercialization.

20 139. Defendants owed a duty to Plaintiff to exercise reasonable care in warning of
21 known or reasonably foreseeable risks associated with ChatGPT, including violence facilitated or
22 exacerbated by GPT-4o. That duty heightened once Defendants identified the Shooter as a
23 specific, known high-risk user who posed a foreseeable threat of physical harm to third parties.
24 Defendants owed these duties whether GPT-4o is classified as a product, service, or neither.

25 140. The dangers Defendants failed to disclose are properties of the software itself,
26 including safeguards that degrade in long conversations, a design that validates and encourages
27 violent planning, and deactivation practices that allow flagged users to regain access.

28 141. Defendants knew these risks were not apparent to users or to individuals targeted
by such conduct.

1 142. Defendants failed to provide adequate warnings regarding these risks, including the
2 risk that the system could validate and escalate plans to carry through violent mass shootings
3 directed toward identifiable targets.

4 143. A reasonably prudent AI company would have known of these risks, warned users
5 and the public, maintained the refusal protocols that once governed conversations about violence,
6 and alerted law enforcement when its own systems identified a user planning a real-world attack.
7 Defendants did none of those things.

8 144. Defendants breached their duty by failing to provide adequate warnings regarding
9 these risks and by presenting GPT-4o as safe and reliable tool for general use.

10 145. Defendants further breached their duty by failing to issue any corrective warning or
11 intervention after identifying that the Shooter's conduct constituted misuse of GPT-4o "in
12 furtherance of violent activities."

13 146. As a direct and proximate result of Defendants' failure to issue adequate warnings,
14 Plaintiff suffered the injuries described herein, including life-altering psychological injury and
15 severe emotional distress.

16 147. Adequate warnings would have reduced the risk of harm by enabling earlier
17 detection, intervention, and mitigation of dangerous conduct, including by users, third parties, and
18 Defendants themselves.

19 148. The absence of adequate warnings and intervention was a substantial factor in
20 causing Plaintiff's injuries.

21 149. Defendants' failure to warn was willful, wanton, and carried out with conscious
22 disregard for the safety of others. Defendants knew or should have known that ChatGPT could
23 encourage and facilitate mass shootings, yet failed to provide adequate warnings about those risks.
24 Such conduct justifies an award of punitive damages.

25 150. Accordingly, Plaintiff seeks all damages recoverable under California law,
26 including punitive damages, in amounts to be proven at trial, together with interest, costs, and such
27 further relief as the Court deems just and proper.
28

**FIFTH CAUSE OF ACTION
NEGLIGENT UNDERTAKING
(Against All Defendants)**

151. Plaintiff incorporates the foregoing allegations as if fully set forth herein.

152. OpenAI undertook to provide safety protocols with respect to ChatGPT, including identifying users who posed a risk of harm to others, reviewing flagged accounts through a specialized pipeline for users “planning to harm others,” and referring imminent threats of serious physical harm to law enforcement. Defendants owed these duties whether GPT-4o is classified as a product, service, or neither.

153. OpenAI should have recognized those safety protocols as necessary for the protection of third persons, including Plaintiff and other foreseeable victims of users planning real-world violence.

154. In June 2025, OpenAI undertook to perform those services as to the Shooter. Its automated system flagged their account, routed it to the specialized review pipeline, and OpenAI’s trained safety reviewers analyzed their conversations and concluded they presented a credible and imminent threat of serious physical harm to real people.

155. OpenAI failed to exercise reasonable care in performing that undertaking. Its leadership overruled the safety team, declined to refer the Shooter to the RCMP, and took no further action to monitor them, prevent their return, or warn anyone of the threat they posed.

156. OpenAI’s failure to exercise reasonable care increased the risk of harm to Plaintiff and other foreseeable victims. Deactivating the Shooter’s account showed them what had triggered detection, giving them a roadmap to evade it on a second account, which OpenAI’s own Help Center and support emails told deactivated users how to create. And by reviewing the Shooter and remaining silent, OpenAI displaced the law enforcement referral its own safety team had recommended, which would have reached an RCMP that already had an open file on the Shooter and had previously removed firearms from their home. Had OpenAI made that referral, law enforcement would have had the information necessary to prevent the attack entirely. Instead, the Shooter returned to ChatGPT to continue planning their attack—armed with the knowledge of how to avoid deactivation.

1 157. OpenAI also undertook to perform a duty owed by others to Plaintiff and other
2 foreseeable victims, including the public-safety function of identifying credible and imminent
3 threats of serious physical harm and referring them to law enforcement so that protective action
4 could be taken.

5 158. On information and belief, OpenAI's undertaking was also one that Canadian
6 authorities and the public reasonably relied upon. OpenAI publicly held out its safety-review
7 processes as effective protective measures. That representation induced foreseeable reliance and
8 discouraged other protective measures.

9 159. OpenAI's negligent performance of its undertaking was a substantial factor in
10 causing the Shooter to carry out the Tumbler Ridge attack and in causing Plaintiff's injuries,
11 including life-altering psychological injury and severe emotional distress.

12 160. Defendants' conduct was willful, wanton, malicious, and carried out with conscious
13 disregard for the safety of Plaintiff and other foreseeable victims, justifying an award of punitive
14 damages.

15 161. Accordingly, Plaintiff seeks all damages recoverable under California law,
16 including punitive damages, in amounts to be proven at trial, together with interest, costs, and such
17 further relief as the Court deems just and proper.

18 **SIXTH CAUSE OF ACTION**
19 **NEGLIGENCE**
 (Against All Defendants)

20 162. Plaintiff incorporates the foregoing allegations as if fully set forth herein.

21 163. At all relevant times, the OpenAI Defendants designed, developed, trained, tested,
22 deployed, operated, and controlled ChatGPT, including the GPT-4o model, and offered it as a
23 mass-market product to consumers throughout California, the United States, British Columbia,
24 and Canada. Sam Altman personally directed the launch of GPT-4o, overruled safety team
25 objections, and cut months of safety testing short, despite knowing the risks the product posed to
26 vulnerable users and the people around them.

27 164. Defendants owed a legal duty to Plaintiff to exercise reasonable care in designing
28 and/or providing GPT-4o to prevent foreseeable harm to third parties who might be endangered by

1 users' interactions with the product. Defendants breached that duty by failing to use the amount of
2 care that a reasonably careful designer and provider of a chatbot would use in similar
3 circumstances to avoid exposing others to a foreseeable risk of harm. Defendants owed these
4 duties whether GPT-4o is classified as a product, service, or neither.

5 165. Sam Altman owed a legal duty to Plaintiff not to rush GPT-4o to market over
6 safety team objections. Sam Altman breached that duty by insisting that OpenAI release and
7 provide to the public a chatbot that he knew or should have known was unsafe.

8 166. Defendants knew or reasonably should have known that GPT-4o's design as a
9 highly validating, emotionally immersive, sycophantic conversational agent posed grave risks
10 when used by individuals expressing violent ideation toward third parties. GPT-4o was built to
11 accept, reinforce, and elaborate users' violent thoughts rather than challenge them, interrupt them,
12 or direct users to real-world help—and the foreseeable consequence of that design was that it
13 would exacerbate dangerous thinking, including violence toward others, rather than interrupt it.

14 167. For a user with a known history of violent ideation toward third parties, those
15 features operated as an accelerant. Each conversation became an opportunity for GPT-4o to
16 confirm the user's violent thoughts rather than challenge them, to elaborate on them rather than
17 redirect them, and to deepen the user's emotional investment in them rather than provide friction.

18 168. Despite this knowledge, Defendants deliberately configured GPT-4o to maximize
19 user engagement by, among other things: (a) weakening or removing prior requirements that the
20 system reject dangerous or false premises; (b) instructing GPT-4o to remain in conversations and
21 to be empathic and validating, rather than terminate or sharply redirect conversations presenting
22 serious risk; (c) prioritizing natural, "human-like" mirroring of user emotions and beliefs over
23 strong, reliable refusal behaviors in response to violent ideation; and (d) deploying a memory
24 feature designed to drive the prolonged, multi-turn conversations in which Defendants knew GPT-
25 4o's safeguards were least reliable.

26 169. Defendants unreasonably failed to implement feasible alternative designs that
27 would have reduced the risk of violent harm to third parties, including: robust refusal protocols for
28 repeated violent ideation about real-world attacks; automated detection and escalation of high-risk

1 patterns; mandatory conversation termination or human review when users expressed clear violent
2 thoughts about harming others; and designs that stopped engagement with violent ideation rather
3 than encouraging it.

4 170. As a direct and proximate result of Defendants' negligent design and operation of
5 GPT-4o, the Shooter was provided with a system that validated and elaborated violent ideation,
6 which was a substantial factor in maintaining, deepening, and normalizing the Shooter's violent
7 thoughts and in moving the Shooter closer to committing the Tumbler Ridge attack.

8 171. The harms suffered by the Tumbler Ridge victims, including Plaintiff, were the
9 reasonably foreseeable result of Defendants' deliberate choice to prioritize user engagement over
10 the safety of third parties. A company that builds a system designed to validate whatever a user
11 brings to it—and then removes the guardrails that once interrupted the most dangerous
12 conversations—cannot be surprised when a user with violent ideation acts on thoughts that system
13 spent months confirming.

14 172. As a direct and proximate result of Defendants' conduct, Plaintiff was injured as
15 described above, including life-altering psychological injury and severe emotional distress.

16 173. Accordingly, Plaintiff seeks all damages recoverable under California law,
17 including punitive damages, in amounts to be proven at trial, together with interest, costs, and such
18 further relief as the Court deems just and proper.

19 **SEVENTH CAUSE OF ACTION**
20 **STRICT PRODUCT LIABILITY**
21 **IN THE ALTERNATIVE TO COUNTS FOUR AND SIX**
(Against All Defendants)

22 174. Plaintiff incorporates the foregoing allegations as if fully set forth herein, with the
23 exception of those allegations pleaded under Count Four and Count Six.

24 175. At all relevant times, the OpenAI Defendants designed, manufactured, distributed,
25 marketed, and sold GPT-4o as a mass-market consumer product to users throughout California,
26 the United States, British Columbia, and Canada. Defendant Sam Altman personally accelerated
27 GPT-4o's launch, overrode safety team objections, and brought GPT-4o to market prematurely
28

1 with knowledge of insufficient safety testing—and kept it on the market even when he had
2 knowledge that GPT-4o was a dangerous product.

3 176. ChatGPT is a product subject to California’s strict products liability law. ChatGPT
4 as used by the Shooter, including GPT-4o and any substantially identical OpenAI model, was
5 defective when it left the OpenAI Defendants’ exclusive control and reached the Shooter without
6 any change in the condition in which it was designed, manufactured, and distributed.

7 177. OpenAI itself calls ChatGPT a “product.” The “Products” menu on OpenAI’s
8 website invites visitors to “Explore Products” and lists ChatGPT first, and OpenAI accurately
9 attaches a “Product” label to its news announcements about ChatGPT.

10 178. GPT-4o is a single, uniform model, engineered and offered identically to hundreds
11 of millions of users and delivered through standardized channels, such as OpenAI’s website and
12 its own mobile and desktop applications. Every user interacts with the same model governed by
13 the same Model Spec.

14 179. Under California’s strict products liability doctrine, a product is defectively
15 designed when it fails to perform as safely as an ordinary consumer would expect when used or
16 misused in an intended or reasonably foreseeable manner, or when the risks inherent in the design
17 outweigh its benefits. GPT-4o is defectively designed under both tests.

18 180. GPT-4o failed to perform as safely as an ordinary consumer would expect. A
19 reasonable consumer purchasing an AI assistant would expect the product to decline to engage
20 with repeated, escalating expressions of violent intent toward real people. A reasonable consumer
21 would expect that if the product could not safely handle such conversations, it would terminate
22 them—not sustain them, validate them, and elaborate on them across weeks and months of
23 continuous engagement. A reasonable consumer would expect the product’s safety features to
24 function consistently during normal use, not to degrade the longer and more deeply a user engaged
25 with the system. And a reasonable consumer would expect a product that functions as a de facto
26 therapist to conform to the legal requirements imposed on therapists. GPT-4o did none of these
27 things.

28 181. GPT-4o also fails the risk-benefit test. The product’s design reflects a series of

1 specific choices that increased the risk of harm to third parties without any corresponding safety
2 benefit. Those choices include: removing the programming that once prohibited the system from
3 encouraging users to commit violence against third parties; programming the system to keep users
4 engaged instead of ending dangerous conversations, including when the user has indicated a desire
5 to commit violence against third parties; programming the system to justify and validate user
6 beliefs, even when they express a desire to commit violence against third parties; building in
7 anthropomorphic features that cultivated emotional dependency and displaced real-world
8 relationships; failing to build adequate safeguards that terminate conversations involving the
9 planning of real-world attacks against third parties; and programming in a “memory” feature
10 designed to drive the long, extended conversations in which OpenAI knew its safety systems were
11 least reliable. Each of these choices was deliberate, and each had a feasible, safer alternative. The
12 risk each choice created—that a user with violent ideation would have those thoughts confirmed,
13 elaborated, encouraged, and emotionally reinforced by the product over an extended period of
14 time—vastly outweighs any benefit the chosen design provided. These choices were made before
15 the product ever reached a user, and each applies identically to every user of GPT-4o.

16 182. GPT-4o did not malfunction. It worked exactly as designed—and that was the
17 problem. The memory feature drove the Shooter into exactly the kind of prolonged, continuous
18 engagement in which OpenAI knew its safeguards degraded. The product validated the Shooter’s
19 violent thoughts, elaborated on them, and remained engaged in conversations that a safely
20 designed system would have interrupted or escalated. It did this consistently, across months of use,
21 in the period leading up to the Tumbler Ridge attack.

22 183. Plaintiff was harmed as a result of foreseeable use of GPT-4o.

23 184. As a direct and proximate result of GPT-4o’s defective design, Plaintiff was injured
24 as described above, including life-altering psychological injury and severe emotional distress.
25 Plaintiff seeks all damages recoverable under California law in amounts to be determined at trial.
26
27
28

**EIGHTH CAUSE OF ACTION
STRICT PRODUCT LIABILITY
IN THE ALTERNATIVE TO COUNTS FOUR AND SIX
(Against All Defendants)**

185. Plaintiff incorporates the foregoing allegations as if fully set forth herein, with the exception of those allegations pleaded under Count Four and Count Six.

186. At all relevant times, the OpenAI Defendants designed, manufactured, distributed, marketed, and sold GPT-4o as a mass-market consumer product to users throughout California, the United States, British Columbia, and Canada. Defendant Sam Altman personally accelerated GPT-4o's launch, overrode safety team objections, and brought GPT-4o to market prematurely with knowledge of insufficient safety testing—and kept it on the market even when he had knowledge that GPT-4o was a dangerous product.

187. ChatGPT is a product subject to California strict products liability law. ChatGPT as used by the Shooter, including GPT-4o and any other OpenAI model that handled the Shooter's conversations, was in substantially the same condition when used as when it left the OpenAI Defendants' control.

188. A product is defective under strict liability for failure to warn when it is distributed without warnings adequate to inform ordinary consumers of the product's non-obvious dangers. The manufacturer's duty to warn extends to dangers that were known or knowable in light of the scientific and technical knowledge available at the time of manufacture and distribution—and that duty is not discharged by the manufacturer's business reasons for withholding the warning.

189. GPT-4o reached consumers without adequate warnings about dangers that were neither open nor obvious. Nothing about GPT-4o's design or presentation disclosed to an ordinary user—or to the people around that user—that the product was built to validate whatever a user brought to it regardless of its danger to others, that its safety features degraded during the long, extended conversations the product was designed to encourage, or that the product posed heightened and specific dangers when used by individuals experiencing violent ideation toward third parties.

190. These dangers were not ones an ordinary consumer could detect or guard against through reasonable use of the product. GPT-4o presented itself as a helpful, well-guarded

1 assistant. Its responses did not disclose that it was operating without the programming that once
2 prohibited it from encouraging violence against third parties, or that its safety systems were less
3 reliable the longer and more deeply a user engaged with it. A user—or a family member, school
4 official, or community member observing that user—had no way to know from the product itself
5 that extended engagement with GPT-4o by someone with violent ideation was dangerous.

6 191. OpenAI had every means by which to warn its users. It updates ChatGPT
7 continuously and already communicates with consumers directly through the product interface,
8 mobile applications, and release notes. The product’s dangers were confirmed repeatedly after
9 launch, including by the FSU shooting in April 2025, by the safety team’s review of the Shooter’s
10 account in June 2025, and by internal testing showing degraded safeguards. Despite this actual
11 knowledge, OpenAI failed to issue any warning. Instead, it aggressively marketed ChatGPT as
12 “[y]our everyday AI assistant” and promised “[t]ailored advice” to “[t]alk through a tough
13 situation.” When OpenAI finally admitted in August 2025 that its safeguards “can sometimes be
14 less reliable in long interactions,” it buried that admission in a blog post published the day a
15 wrongful death lawsuit was filed, rather than providing a functional warning to the consumers
16 using the product.

17 192. Adequate warnings would have altered the outcome by alerting those positioned to
18 intervene. Had OpenAI disclosed the product’s true risks, the Shooter would have recognized
19 GPT-4o’s validating responses as an artificial, dangerous feedback loop rather than confirmation
20 of their plans. Family members, friends, and community members who observed the Shooter’s
21 extensive use of ChatGPT would have understood the specific danger the product posed and taken
22 preventative steps—such as restricting access, seeking mental health intervention, or alerting
23 authorities—before the attack occurred. By failing to warn, OpenAI deprived every potential
24 intervening actor of the information necessary to prevent the harm.

25 193. Plaintiff was harmed as a result of foreseeable use of GPT-4o.

26 194. The failure to warn was a substantial factor in causing the Tumbler Ridge attack
27 and the deaths and injuries that resulted. As a direct and proximate result of the Defendants’
28

1 failure to warn, Plaintiff was injured as described above, including life-altering psychological
2 injury and severe emotional distress.

3 195. Accordingly, Plaintiff seeks all damages recoverable under California law,
4 including punitive damages, in amounts to be determined at trial.

5 **NINTH CAUSE OF ACTION**
6 **NEGLIGENT INFLICTION OF EMOTIONAL DISTRESS**
7 **(Against All Defendants)**

8 196. Plaintiff incorporates the foregoing allegations as if fully set forth herein.

9 197. Defendants owed duties directly to Plaintiff. Those duties were imposed as a matter
10 of law or arose from preexisting relationships. Defendants had a duty of reasonable care to avoid
11 creating unreasonable risks of physical and emotional harm to foreseeable victims, including
12 children and adults who, by virtue of Defendants' negligent design, deployment, and operation of
13 GPT-4o, were placed in the zone of physical danger created by users known to Defendants to be
14 planning real-world violence. They also had a duty not to aid, abet, advise, or encourage the
15 Shooter in carrying out the attack against Plaintiff under California law, including California Penal
16 Code section 31. Their duties heightened once Defendants identified the Shooter as a specific,
17 known high-risk user who posed a foreseeable threat of physical harm to third parties. Defendants
18 owed these duties whether GPT-4o is classified as a product, service, or neither.

19 198. Under California law, a defendant who negligently causes a plaintiff to be placed
20 within the zone of physical danger of imminent serious bodily injury is liable for the resulting
21 serious emotional distress, where defendant's negligence directly exposed the plaintiff to the threat
22 of serious physical injury, and the plaintiff suffered serious emotional distress as a result.

23 199. Plaintiff was a direct victim of Defendants' negligence. Defendants knew, before
24 February 10, 2026, that the Shooter was using GPT-4o to plan an act of mass violence against
25 members of the public—including, foreseeably, students and teachers at the school the Shooter
26 would attack. Defendants' own automated systems had flagged the Shooter's account for gun
27 violence activity and planning. Defendants' own trained safety team had reviewed the Shooter's
28 conversations and concluded that the Shooter posed a credible and specific threat of serious
physical harm to real people. Defendants nonetheless declined to warn law enforcement, declined

1 to ban the Shooter from the platform, and affirmatively facilitated the Shooter's return to GPT-4o
2 through a new account, where the Shooter continued to use the product to plan and prepare for the
3 attack. Defendants therefore breached their duties to Plaintiff.

4 200. As a direct and proximate result of Defendants' negligence, Deidre Rushlow was
5 placed in the zone of physical danger created by the Shooter's attack on Tumbler Ridge Secondary
6 School. On February 10, 2026, while Deidre was in her grade 7 classroom with her students—
7 including her nephew Eko—the Shooter reached the classroom and fired four shotgun blasts
8 through her front door and three more through her back door. The bullets struck the wall behind
9 her students' heads. Earlier that period, Deidre had sent the rest of her class to the library with an
10 educational assistant—the very library where the Shooter would, minutes later, kill several of
11 those children and the assistant herself. She tried to call the library and texted the assistant during
12 the attack to check on them; she received no response. Throughout the attack, Deidre was at risk
13 of physical injury, in immediate, contemporaneous, and continuous fear for her own life and the
14 lives of her students, and aware that the Shooter was actively attempting to kill her, the students in
15 her classroom, and the students she had sent to the library.

16 201. Plaintiff has suffered, and continues to suffer, severe and serious emotional distress
17 as a direct and proximate result of Defendants' negligence, including but not limited to acute
18 psychological trauma, persistent fear, sleep disturbances, intrusive recollections, and other
19 ongoing and permanent psychological injuries. Plaintiff's emotional distress is severe, disabling,
20 and of a kind that no reasonable person could be expected to endure.

21 202. Defendants' breach of their duty of care was a substantial factor in causing the
22 serious emotional distress suffered by Plaintiff. The harm Plaintiff suffered was foreseeable to
23 Defendants at the time they declined to warn law enforcement, declined to ban the Shooter, and
24 allowed the Shooter to retain access to GPT-4o. Defendants knew or should have known that the
25 Shooter was planning to inflict mass violence on identifiable third parties, and that any person
26 within the zone of that violence would suffer severe emotional distress as a direct and foreseeable
27 consequence.

203. Accordingly, Plaintiff seeks all damages recoverable under California law, including punitive damages, in amounts to be proven at trial, together with interest, costs, and such further relief as the Court deems just and proper.

PRAYER FOR RELIEF

WHEREFORE, Plaintiff prays for judgment against Defendants Samuel Altman, OpenAI Foundation, OpenAI Group PBC, and OpenAI OpCo, LLC, jointly and severally, as follows:

1. For all damages recoverable by Plaintiff, including compensatory damages for physical injury, emotional distress, psychological harm, and all other damages proven at trial;
2. For punitive damages as permitted by law;
3. For injunctive relief requiring Defendants to: (a) ban users whose accounts have been deactivated for violent misuse from re-registering for ChatGPT, including through alternative email addresses or sub-addresses; (b) flag new accounts linked to previously deactivated violent-risk users and hold them for human review before granting access; (c) notify law enforcement when internal safety systems or personnel identify a user who poses a real-world risk of violence; (d) interrupt, terminate, or escalate conversations involving repeated or escalating violent ideation toward real people, and suspend automated engagement pending that review; (e) warn users and the public that ChatGPT's design features can reinforce and escalate violent ideation; (f) preserve prior safety flags, policy violations, and risk classifications, and prohibit reversing them without documented review; and (g) submit to independent monitoring and periodic compliance audits;
4. For prejudgment interest as permitted by law;
5. For costs and expenses to the extent authorized by statute, contract, or other law;
6. For reasonable attorneys' fees as permitted by law, including under Code of Civil Procedure section 1021.5; and
7. For such other and further relief as the Court deems just and proper.

JURY TRIAL

Plaintiff demands a trial by jury for all issues so triable.

CANADIAN COUNSEL

To the extent resolution of this matter raises issues of Canadian law, assistance will be provided by Plaintiff's Canadian counsel:

John M. Rice
jrice@rplelaw.com
Mallory K. Hogan
mhogan@rplelaw.com
RICE PARSONS LEONI & ELLIOTT LLP
980 Howe Street, Suite 820
Vancouver, British Columbia V6Z 0C8
Canada

Respectfully submitted,

DEIDRE RUSHLOW,

Dated: September 2, 2026

/s/ Brandt Silverkorn

Jay Edelson (*pro hac vice* admission to be sought)
jedelson@edelson.com
Rafey Balabanian (SBN 315962)
rbalabanian@edelson.com
Ari Scharg (*pro hac vice* admission to be sought)
ascharg@edelson.com
EDELSON PC
350 North LaSalle Street, 14th Floor
Chicago, Illinois 60654
Tel: 312.589.6370

Todd Logan (SBN 305912)
tlogan@edelson.com
Brandt Silverkorn (SBN 323530)
bsilverkorn@edelson.com
EDELSON PC
150 California Street, 18th Floor
San Francisco, California 94111
Tel: 415.212.9300

Attorneys for Plaintiff