

**UNITED STATES DISTRICT COURT
NORTHERN DISTRICT OF ILLINOIS
EASTERN DIVISION**

JESSICA BERLINER, derivatively on behalf
of NVIDIA CORPORATION,

Plaintiff,

vs.

JEN-HSUN HUANG, STEPHEN C. NEAL,
TENCH COXE, JOHN O. DABIRI, DAWN
HUDSON, HARVEY C. JONES, MELISSA
B. LORA, A. BROOKE SEAWELL, AARTI
SHAH, MARK A. STEVENS, COLETTE M.
KRESS, AJAY K. PURI, DEBORAH
SHOQUIST, and TIMOTHY S. TETER,

Defendants,

– and –

NVIDIA CORPORATION,

Nominal Defendant.

No.

**VERIFIED STOCKHOLDER
DERIVATIVE COMPLAINT**

JURY TRIAL DEMANDED

VERIFIED STOCKHOLDER DERIVATIVE COMPLAINT

Plaintiff Jessica Berliner (“Plaintiff”), derivatively on behalf of NVIDIA Corporation (“NVIDIA” or the “Company”), alleges the following based on personal knowledge as to Plaintiff and Plaintiff’s acts and as to all other matters on information and belief based on the investigation of Plaintiff’s counsel, which included, among other things, a review of legal and regulatory filings, press releases, NVIDIA’s online documents, media reports about NVIDIA and other public statements issued by the Company. Plaintiff believes that substantial additional evidentiary support will exist for the allegations set forth herein after a reasonable opportunity for discovery.

I. NATURE OF THE ACTION

1. This is a stockholder derivative action brought by Plaintiff on behalf of NVIDIA against certain of its officers and directors for adopting and implementing an unlawful business strategy whereby NVIDIA used copyrighted and unlicensed materials, including books, videos, and commercial voices, to develop its Artificial Intelligence (“AI”) services.

2. NVIDIA is a diversified technology company founded in 1993 that originally focused on computer-graphics hardware, Graphics Processing Units (“GPUs”), but has since expanded out to other computationally intensive fields, including software such as NVIDIA’s “Compute Unified Device Architecture,” and hardware such as NVLink/NVLink Switch, for training and operating AI software programs. NVIDIA’s hardware and software are used by all “Frontier AI” companies that develop the most advanced AI systems, which has resulted in NVIDIA becoming a valuable worldwide company. Frontier AI typically refers to the most advanced, general purpose AI models, large-scale systems that sit at the “frontier” of current capabilities in areas such as reasoning, multimodal understanding, and autonomous task execution.¹

3. NVIDIA AI models include multiple AI software programs called large language models (“LLM’s”) created, maintained, and commercialized by NVIDIA. NVIDIA’s LLM’s include its primary NeMo Megatron models, Nemo Megatron-GPT 20B, Nemo Megatron-GPT 1.3B, Nemo Megatron-GPT 5B, and Nemo Megatron-TF 3B. Each of these Nemo Megatron LLM models were hosted on the Hugging Face website, trained on “The Pile,” prepared by EleutherAI, which used the Books3 pirated dataset,

¹ Dana Larson, *Understanding Frontier AI: Definitions, Models, and Business Impact*, CROWDSTRIKE (May 26, 2026).

which included approximately 196,640 pirated books.

4. NVIDIA unlawfully copied these copyrighted works multiple times to train its language models, including from multiple known pirated “shadow libraries,” The Pile, Bibliotek, and Anna’s Archive. NVIDIA also downloaded the SlimPajama data, which was created by cleaning and duplicating the RedPajama dataset, and included the Books3 dataset. In addition, NVIDIA also downloaded copyrighted materials based on shadow libraries, including LibGen, Sci-Hub, and Z-Library.

5. NVIDIA AI models and platforms also include Vision-Language Models (“VLMs”), NVIDIA AI Video, a large-scale generative text-to-video AI model named Cosmos. NVIDIA developed datasets for Cosmos by collaborating and communicating via chats on the “Slack” platform, including a Slack channel entitled “#cosmos-dataset-creation,” and infringed copyright materials through the use of the HD-VG-130M, HDVILA100M, and HowTo100M datasets, comprised of YouTube videos that NVIDIA extracted from YouTube through circumvention of YouTube’s technological protection measures (“TPMs”). HD-VG-130M was comprised of 1,549,408 YouTube videos, HDVILA100M was comprised of 3,098,462 YouTube videos, and HowTo100M was comprised of 1,549,408 YouTube videos, all of which were broken into video clips.

6. NVIDIA AI models also include the training and deployment of generative AI models, voice synthesis, voice cloning, and voice-conversion models, that developed multiple commercial voice models (“CVMs”) for NVIDIA’s commercial voice AI products. NVIDIA distributes these CVMs through NVIDIA AI Enterprise, the NVIDIA API Catalog, NVIDIA NiM microservices, and open-weight models releases on Hugging Face. These CVMs include Magpie TTS Zeroshot, Magpie TTS Flow, FUGATTO, PersonaPlex, Nemotron 4 VoiceChat, and the Canary and Parakeet speech models.

NVIDIA ingested within each of these CVMs hundreds of thousands of hours of human speech recordings and extracted the unique biometric signatures and the voiceprints of the speakers from those recordings. In violation of the Illinois Biometric Information Privacy Act (“BIPA”), the Defendants failed to identify the source speakers, provide written notice of the specific purpose and duration of collection, and obtain a written release from each speaker before ingesting that speaker’s recording into the training pipeline, which resulted in damages to the source speakers.

7. Defendants are all directors and/or officers of NVIDIA who knew of the copyright and BIPA violations being perpetrated through NVIDIA’s AI models, including LLMs, VLMs, and CVMs. Rather than utilize a clean dataset that did not include copyrighted books, YouTube videos, and commercialized voiceprints, Defendants followed the “ask forgiveness not approval” model and adopted and implemented an unlawful plan to use a pirated dataset to develop the Company’s AI services. Defendants’ misconduct constitutes bad faith and a violation of privacy law in directing, approving, or permitting the misconduct.

8. Predictably, Defendants’ misconduct has resulted in multiple copyright holders filing lawsuits against NVIDIA based on NVIDIA’s failure to compensate them for downloading, copying, storing, or using their copyrighted works, YouTube videos, and commercialized voiceprints, for which NVIDIA is now facing potential massive liability.

9. Defendants were well aware of all these issues with NVIDIA’s downloading, copying, and storing, and using the copyright holders’ copyrighted works and YouTube videos, and commercialized voiceprints without their permission or compensation.

10. Defendants were also well aware of the potential liability for NVIDIA’s

unauthorized utilization of copyright infringement based on multiple individual and class action lawsuits against other corporations for copyright infringements relating to AI platforms, including: *Kadrey v. Meta Platforms, Inc.*, No. 3:23-cv-3417-VC (N.D. Cal. 2023); *Concord Music Group Inc. v. Anthropic PBC*, No. 5:24-cv-03811-EKL (N.D. Cal. 2024) (“*Concord I*”); *Bartz v. Anthropic PBC*, No. 4:24-cv-05417-AMO (N.D. Cal. 2024); *In re Mosaic LLM Litigation*, No. 3:24-cv-01451-CRB (N.D. Cal. 2024); *Hendrix v. Apple Inc.*, No. 4:25-cv-07558-YGR (N.D. Cal. 2025); *Lyon v. Adobe Inc.*, No. 3:25-cv-10732-JSC (N.D. Cal. 2025); *James v. Snowflake Inc.*, No. 2:25-cv-00108-BMM (D. Mont. 2025); *Tanzer v. Salesforce Inc.*, No. 3:25-cv-08862-CRB (N.D. Cal. 2025); *Bird v. Microsoft Corp.*, No. 1:25-cv-05282 (S.D.N.Y. 2025); *Ted Entertainment, Inc. v. OpenAI, Inc.*, No. 5:26-cv-02935-BLF (N.D. Cal. 2026); *Kleiner v. Adobe Inc.*, No. 3:26-cv-01218-JSC (N.D. Cal. 2026); and *Ted Entertainment, Inc. v. Amazon.com, Inc.*, 2:26-cv-01134 (W.D. Wash. 2026).

11. Defendants are also well aware of multiple, similar BIPA lawsuits filed against other AI platform development corporations, including: *Basich v. Microsoft Corp.*, No. 2:26-cv-00422 (W.D. Wash. 2026); *Amer v. Eleven Labs Inc.*, No. 1:26-cv-05437 (N.D. Ill. 2026); *Marin v. Alphabet, Inc.*, No. 1:26-cv-05436 (N.D. Ill. 2026); *Marin v. Meta Platforms, Inc.*, No. 1:26-cv-05438 (N.D. Ill. 2026); *Delgado v. Meta Platforms, Inc.*, No. 3:23-cv-04181-SI (N.D. Cal. 2023); *Lacour v. Apple Inc.*, No. 1:26-cv-05536 (N.D. Ill. 2026); *Dorcus v. Adobe Inc.*, No. 1:26-cv-05575 (N.D. Ill. 2026), and *Flowers v. Microsoft Corp.*, No. 1:26-cv-05491 (N.D. Ill. 2026).

12. Indeed, another AI development company, Anthropic, paid damages of \$1.5 billion to resolve the *Concord I* action and is now subject to additional damages of up to \$3 billion in a second action filed in 2026, *Concord Music Group Inc. v. Anthropic PBC*,

No. 5:26-cv-00880-EKL (N.D. Cal. 2026) (“*Concord II*”). This tidal wave of litigation filed against AI developers like NVIDIA served as blazing red flags alerting Defendants of the need to prevent NVIDIA from committing and continuing to commit the same unlawful conduct.

13. Copyright holders and commercial voiceprint holders have filed three class actions against NVIDIA for copyright infringement and BIPA violations. On March 9, 2024, copyrighted book authors filed their class action lawsuit against NVIDIA, *Nazemian v. NVIDIA Corp.*, No. 4:24-cv-01454-JST (N.D. Cal.). On November 26, 2025, intellectual property video creators filed their initial class action lawsuit against NVIDIA, *Ted Entertainment, Inc. v. NVIDIA Corp.*, No. 5:25-cv-10287-EJD (N.D. Cal.). On May 12, 2026, commercial voice print creators filed their class action lawsuit against NVIDIA, *Rogers v. NVIDIA Corp.*, No. 1:26-cv-05478 (N.D. Ill.). Thus, Defendants were, and are, well aware of the fraudulent and bad faith misconduct that has resulted in the copyright infringement and BIPA class action filings, and the potential damages NVIDIA may have to pay.

14. NVIDIA is now saddled with having to defend itself in copyright infringement cases and BIPA violation cases as a result of Defendants’ misconduct and is facing potentially massive liability and related costs and reputational damages, as well as potential loss of customers and/or claims from NVIDIA customers who unwittingly used the copyright and BIPA infringing NVIDIA services.

II. JURISDICTION AND VENUE

15. This Court has subject matter jurisdiction over this action under 28 U.S.C. § 1331 because of claims arising under 15 U.S.C. §§ 78j(b) and 78n(a), and SEC regulations 10b-5 and 14a-9 promulgated thereunder, 17 C.F.R. §§ 240.10b-5 and 240.14a-

9. This Court has supplemental jurisdiction pursuant to 28 U.S.C. § 1367 as to the state law claims alleged, as they arise out of the same transactions and occurrences as the federal claims.

16. This Court has specific jurisdiction over each named non-resident Defendant because each of the Defendants maintains sufficient minimum contacts in this District to render jurisdiction by this Court permissible under traditional notions of fair play and substantial justice. At the time of the alleged wrongdoing detailed herein, NVIDIA had a principal location in this District, each of the Defendants was a director and/or officer of NVIDIA, and Defendants' conduct was purposefully directed at this District. Exercising jurisdiction over any non-resident Defendant is also reasonable under the circumstances.

17. Venue is proper in this Court pursuant to 28 U.S.C. § 1391(b) and 15 U.S.C. § 78aa. NVIDIA has a principal location in this District and is subject to personal jurisdiction in this District due, in part, to violations of the BIPA based on extraction of voiceprints of Illinois residents. In addition, a substantial portion of the transactions and wrongs complained of herein, including the Defendants' primary participation in the wrongful acts detailed herein in violation of fiduciary duties owed to NVIDIA occurred in this District, and Defendants have received substantial compensation by doing business in this District and engaging in numerous activities that had an effect in this District.

III. THE PARTIES

18. Plaintiff Jessica Berliner is a current stockholder of NVIDIA and has continuously owned shares of NVIDIA common stock at all times relevant herein.

19. Nominal defendant NVIDIA is a Delaware corporation headquartered in Santa Clara, California, with a primary office located in Champaign, Illinois. NVIDIA

common stock trades on the Nasdaq under the ticker symbol “NVDA.”

20. Defendant Jen-Hsun Huang (“Huang”) is on NVIDIA’s board of directors (the “Board”) and has served as NVIDIA’s Chief Executive Officer (“CEO”) and President since he founded NVIDIA in 1993. Huang received roughly \$36.3 million in compensation for 2025, \$49.9 million in compensation for 2024, and \$34.2 million in compensation for 2023, all approved by his co-defendants on the Board. Huang currently holds approximately 3.5% of NVIDIA’s shares, a total of 851,983,603 shares, worth over \$173.9 billion.

21. Defendant Stephen C. Neal (“Neal”) is the Lead Director of the Board and has been a director since 2019. Neal is the Chairman Emeritus & Senior Counsel at Cooley LLP. Neal received \$363,809 in Board-related compensation for 2025.

22. Defendant Tench Coxe (“Coxe”) has been a director since 1993 and is a member of the Audit Committee. Coxe is a former Managing Director at Sutter Hill Ventures. Coxe received \$363,809 in Board-related compensation for 2025. Coxe currently holds 30,555,240 NVIDIA shares.

23. Defendant John O. Dabiri (“Dabiri”) has been a director since 2020. Dabiri is the Centennial Professor of Aeronautics and Mechanical Engineering at the California Institute of Technology. Dabiri received \$363,809 in Board-related compensation for 2025.

24. Defendant Dawn Hudson (“Hudson”) has been a director since 2013. Hudson is the former Chief Marketing Officer of the NFL, and the former CEO of Pepsi-Cola North America. Hudson received \$363,809 in Board-related compensation for 2025.

25. Defendant Harvey C. Jones (“Jones”) has been a director since 1993 and is a member of the Audit Committee. Jones is the Managing Partner of Square Wave

Ventures. Jones received \$363,809 in Board-related compensation for 2025. Jones currently holds 7,002,787 NVIDIA shares.

26. Defendant Melissa B. Lora (“Lora”) has been a director since 2023 and is a member of the Audit Committee. Lora is the former President of Taco Bell International. Lora received \$363,809 in Board-related compensation for 2025.

27. Defendant A. Brooke Seawell (“Seawell”) has been a director since 1997 and is the Chairperson of the Audit Committee. Seawell is the Venture Partner of New Enterprise Associates. Seawell received \$363,809 in Board-related compensation for 2025. Seawell currently holds 2,505,707 NVIDIA shares.

28. Defendant Aarti Shah (“Shah”) has been a director since 2020 and is a member of the Audit Committee. Shah is the former Senior Vice President and Chief Information and Digital Officer of Eli Lilly and Company. Shah received \$363,809 in Board-related compensation for 2025.

29. Defendant Mark A. Stevens (“Stevens”) has been a director since 2008 and previously served as a director from 1993 to 2006. Stevens is the Managing Partner of S-Cubed Capital. Stevens received \$363,809 in Board-related compensation for 2025. Stevens currently holds 34,068,672 shares of NVIDIA.

30. Defendant Colette M. Kress (“Kress”) is one of NVIDIA’s Executive Vice Presidents (“EVP”) and its Chief Financial Officer (“CFO”). Kress oversees NVIDIA’s accounting, business support, financial planning and analysis, treasury, investor relations, internal audit, and tax functions. Kress received roughly \$14.4 million in compensation for 2025, \$21.4 million in compensation for 2024, and \$13.3 million in compensation for 2023, all approved by her co-defendants on the Board.

31. Defendant Ajay K. Puri (“Puri”) is NVIDIA’s EVP of Worldwide Field

Operations, CFO. Puri received \$14.8 million in 2025 compensation, \$21.6 million in 2024 compensation, and \$13.6 million for 2023 compensation, all approved by the Board.

32. Defendant Deborah Shoquist (“Shoquist”) is NVIDIA’s EVP of Operations. Shoquist received approximately \$14.3 million in 2025 compensation, \$19.2 million in 2024 compensation, and \$11.1 million for 2023 compensation, all approved by the Board.

33. Defendant Timothy S. Teter (“Teter”) is one of NVIDIA’s EVPs, the General Counsel and Secretary. Teter received roughly \$14.3 million in compensation for 2025, \$19.2 million in compensation for 2024, and \$11.1 million in compensation for 2023, all approved by his co-defendants on the Board.

34. Defendants Huang, Coxe, Dabiri, Hudson, Jones, Lora, Neal, Seawell, Shah, and Stevens are collectively referred to herein as the “Director Defendants.” Defendants Coxe, Jones, Lora, Seawell, and Shah are collectively referred to herein as the “Audit Committee Defendants.” Defendants Huang, Kress, Puri, Shoquist, and Teter are collectively referred to herein as the “Officer Defendants.” The Director Defendants and the Officer Defendants are collectively referred to herein as the “Defendants.”

IV. DUTIES OF THE DEFENDANTS

A. General Fiduciary Duties of Officers and Directors

35. By reason of their positions as officers and directors of the Company, each of the Defendants owed and owe NVIDIA and its stockholders fiduciary obligations of care and loyalty and were and are required to use their utmost ability to control and manage NVIDIA in a fair, just, honest, and equitable manner. Defendants were and are required to act in furtherance of the best interests of NVIDIA and not in furtherance of their personal interest or benefit.

36. Each Defendant owes to NVIDIA and its stockholders the fiduciary duty to

exercise good faith and diligence in the administration of the Company and in the use and preservation of its property and assets and the highest obligation of fair dealing.

37. Defendants, because of their positions of control and authority as directors and/or officers of NVIDIA, were able to and did, directly and/or indirectly, exercise control over the wrongful acts complained of herein.

38. Defendants' conduct complained of herein involves a knowing and culpable violation of their obligations as directors and/or officers of NVIDIA, the absence of good faith on their part, or a reckless disregard for their duties to the Company and its stockholders that Defendants were aware or should have been aware posed a risk of serious injury to the Company.

39. As senior officers and directors of a publicly-traded company whose common stock was registered with the SEC pursuant to the Securities Exchange Act (the "Exchange Act") and traded on the Nasdaq, the Defendants also had a duty to prevent and not to effect the dissemination of inaccurate and untruthful information with respect to the Company's financial condition, its Code of Conduct, operations, products, management, and internal controls, including the dissemination of false and/or materially misleading information regarding the Company's business, prospects, and operations, and had a duty to cause the Company to disclose in its regulatory filings with the SEC all those facts described in this Complaint that it failed to disclose, so that the market price of the Company's common stock would be based upon truthful, accurate, and fairly presented information.

40. To discharge their duties, the officers and directors of NVIDIA were required to exercise reasonable and prudent supervision over the management, policies, practices, and controls of the financial affairs of the Company.

41. At all times relevant hereto, the Defendants were the agents of each other and of NVIDIA and were at all times acting within the course and scope of such agency.

42. Because of their advisory, executive, managerial, and directorial positions with NVIDIA, each of the Defendants had access to adverse, non-public information about the Company.

43. Defendants, because of their positions of control and authority, were able to and did, directly or indirectly, exercise control over the wrongful acts complained of herein, as well as the contents of the various public statements issued by NVIDIA.

B. NVIDIA's Code of Conduct

44. In addition to the fiduciary obligations owed to NVIDIA, each of the Defendants were also subject to layers of NVIDIA policies prohibiting certain conduct.

45. All Officer Defendants are subject to all facets of the Company's Code of Conduct, and all Director Defendants are subject to the relevant sections of the Company's Code of Conduct as set forth below, which purports to establish the principles of business conduct that NVIDIA considers fundamental in its operations worldwide. In particular, the Code of Conduct provides that each officer, director and employee must protect assets, data, and personal information, and put the Company interests ahead of their own. Specifically, for "Trust," each officer, director and employee must:

- (a) Use NVIDIA assets only for ethical and legal purposes that benefit the Company and its shareholders and should not be used to engage in outside commercial activities or illegal activities, or to create, store, or send content that others might find offensive.
- (b) Keep personal use of NVIDIA's assets to a minimum and avoid any use that might lead to loss or damage.

- (c) Be responsible for all of NVIDIA's assets and intellectual property, and take care to prevent the loss, theft, unauthorized selling or reselling, and misuse of NVIDIA assets, IP, or documents.
- (d) Protect confidential or proprietary information whether it belongs to NVIDIA or others.
- (e) Use patents and copyrighted materials belonging to others only after obtaining the proper licenses and permissions.
- (f) Cite others' trademarks in accordance with the owners' guidelines.
- (g) Don't use NVIDIA networks, computers, or other resources to acquire, share, or store copyrighted material that isn't properly licensed.

46. In order to comply with duties and obligations for "Reporting Concerns," each officer, director and employee must:

- (a) Act when aware of misconduct.
- (b) Report concerns or suspected violations of the Code of Conduct or policies to a manager, a human resources or legal representative, or NVIDIA-Compliance.
- (c) As a member of the Compliance Committee, investigate reports of suspected violations of the Code of Conduct promptly, thoroughly, and in accordance with legal obligations.
- (d) For activities related to open source software, open source software is usually collectively developed software with its

source code made available under an open source license. Before using, modifying, or distributing any open source software for NVIDIA infrastructure, or as part of a NVIDIA product or service development effort, you must receive management and legal approval. For additional information on how to submit requests, visit Open Source at NVIDIA. This website also includes information about personal contributions to open source and required approvals.

47. As discussed below, the Code of Conduct also includes material misrepresentations in concert with the allegations set forth herein:

- (a) “NVIDIA pioneers accelerated and AI computing innovating across the entire stack, the whole data center, and from cloud to robotic systems.” “We prioritize developing trustworthy AI to ensure this powerful technology provides benefits.” “We value diversity and inclusion, which enables us to achieve this goal. Additionally, we recognize our responsibility to maintain high ethical standards and identify concerns”
- (b) “We don’t engage in activities that might limit competition or violate antitrust law.” “We gather competitive information with care, seeking only data that is publicly available or licensed to us.” “We’re committed to full, fair, accurate, timely, and clear disclosures in reports and documents that we file with or submit to government agencies, as well as in other public communications or in material that we develop for internal use.”

- (c) “We support the use of our products for socially beneficial purposes and work to address the risks of potential harm.” “We ensure our customers can trust our products by sharing important information about how they are built and work.” “When it comes to artificial intelligence, we take a proactive approach to identify and address issues of trustworthiness concerns through the development process—from date of collection to decommissioning. When necessary, our AI Ethics Committee considers the implications of operating product.”
- (d) “We don’t create misleading impressions in any advertising, marketing, or sales materials or presentations and don’t make false or illegal claims about competitors or their offerings.”
- (e) “For our artificial intelligence products, we use model cards to explain how a model was designed, trained, and tested, how it performs, and its intended purposes.”

C. NVIDIA’s Audit Committee Charter

48. The Audit Committee Defendants are subject to heightened responsibilities to identify, assess, and mitigate risks associated with the Company’s operations.

49. According to the 2026 Proxy, the Audit Committee:

- (a) Oversees our corporate accounting and financial reporting process;
- (b) Oversees our internal audit function;

- (c) Determines and approves the engagement, compensation, retention, and termination of the independent registered public accounting firm;
- (d) Evaluates the performance and qualifications of our independent registered public accounting firm;
- (e) Reviews and approves the retention of the independent registered public accounting firm for permissible audit and non-audit services;
- (f) Confers with management and our independent registered public accounting firm on the results of the annual audit, our quarterly financial statements and results, and the effectiveness of internal control over financial reporting, including those regarding information security;
- (g) Reviews the financial statements to be included in our quarterly reports on Form 10-Q and annual reports on Form 10-K;
- (h) Reviews earnings press releases and the substance of financial information and outlook provided to investors and analysts on earnings calls;
- (i) Adopts and maintains policies regarding preapproval of employment of individuals employed or formerly employed by auditors and engaged on our account;
- (j) Prepares their report as required to be included by SEC rules in our annual proxy statement or annual report on Form 10-K;

- (k) Establishes procedures for the receipt, retention, and treatment of complaints we receive regarding accounting, internal accounting controls or auditing matters and the confidential and anonymous submission by employees of concerns regarding accounting or auditing matters;
- (l) Oversees risks related to financial reporting and exposures, internal audit functions, regulatory, and accounting policies; and
- (m) Reviews and reports on the adequacy and effectiveness of the Company's information security policies and practices and the internal controls regarding information security risks.

50. Under the Audit Committee Charter, the Audit Committee Defendants:

- (a) [S]hall have full access to all books, records, facilities and personnel of the Company as deemed necessary or appropriate by any member of the Committee to discharge his or her responsibilities hereunder.
- (b) Review and reassess annually the adequacy of this Charter and submit any suggested changes to the Charter to the Board for its consideration.
- (c) Review and recommend to the Board, upon completion of the annual audit, the financial statements, any internal controls report, and as appropriate, "Risk Factors" to be included in the Company's Annual Report on Form 10-K and any internal controls report. Review should include oversight of the Auditors' assessment of the quality, not just acceptability, of

accounting principles, the reasonableness of significant judgments and estimates including material changes in estimates), the nature of significant risks and exposures, any audit adjustments noted or proposed by the Auditors (whether “passed” or implemented in the financial statements), the adequacy of the disclosures in the financial statements and any other matters required to be communicated to the Committee by the Auditors.

- (d) Review, in consultation with management, the Internal Auditors and the Auditors, the Company’s guidelines and policies with respect to risk assessment, risk management, internal control over financial reporting and disclosure controls and procedures. Understand any material changes to policies due to changes in systems, technology, management structure or other processes for disclosure consideration.
- (e) Discuss with management and the Auditors the results of the Auditors’ review of the Company’s quarterly financial statements, prior to public disclosure of quarterly financial information, if practicable, or filing with the SEC of the Company’s Quarterly Report on Form 10-Q, and any other matters required to be communicated to the Committee by the Auditors under applicable standards.
- (f) Review generally with management and the Auditors, as appropriate, earnings press releases, as well as the substance of

financial information and earnings outlook provided to analysts and ratings agencies. The Chairperson of the Committee may represent the entire Committee for purposes of this discussion.

- (g) Review periodically, either individually or as a committee, and discuss with management and the Auditors, as appropriate, the Company's disclosures contained under the caption "Management's Discussion and Analysis of Financial Condition and Results of Operations" in its periodic reports to be filed with the SEC.
- (h) Review with management, the Internal Auditors and the Auditors significant issues that arise regarding accounting principles and financial statement presentation, including the adoption of new, or material changes to existing, critical accounting policies or to the application of those policies; the potential effect of alternative accounting policies available under GAAP and the treatment recommended by the Auditors; the Company's disclosure of non-GAAP financial measures; the potential impact of regulatory and accounting initiatives; any off-balance sheet structures and any other significant reporting issues and judgments.
- (i) Review and discuss with management the Company's capital return program and make recommendations to the Board regarding the Company's capital return program for both dividends and share repurchases.

- (j) Establish procedures for the receipt, retention and treatment of complaints received by the Company regarding accounting, internal accounting controls or auditing matters, including the confidential and anonymous submission by employees of concerns regarding questionable accounting or auditing matters as required by applicable law and as deemed appropriate.
- (k) Investigate any matter brought to the attention of the Committee within the scope of its duties if, in the judgment of the Committee, such investigation is necessary or appropriate. The Chairperson of the Committee may represent the entire Committee in making such determination.
- (l) Review periodically with senior management, in conjunction with the NCGC, any proposed revisions to, and compliance with, the Company's Code of Conduct and Financial Team Code of Conduct.
- (m) Review, at least annually, and approve the Company's decision to enter into swaps or other derivative transactions that are exempt from exchange-execution and clearance under "end-user exception" regulations established by the Commodity Futures Trading Commission, and review and discuss with management the Company's global treasury activities and policies, including the Company's Foreign Exchange Policy and the Company's use of swaps.

- (n) Review and discuss with management and the Auditors the Company's processes and policies on risk identification, management and assessment in all areas of the Company's business, including financial and accounting. Areas of focus shall include the Company's policies and other matters relating to the Company's investments, cash management and foreign exchange management, major financial risk exposures, insurance coverage and strategies, and the steps taken by management to monitor and mitigate or otherwise control these exposures and to identify future risks.
- (o) Assist the Board in its oversight of cybersecurity matters and review the adequacy and effectiveness of the Company's information security policies and practices and the internal controls regarding information security risks.
- (p) Review and discuss with management the Company's global tax principles and strategy, tax planning activities, as well as any significant tax developments affecting the Company, and related risks.
- (q) Prepare annually a report to stockholders as required by the SEC and to be included in the Company's proxy statement.
- (r) Report to the Board with respect to material issues that arise regarding the quality or integrity of the Company's financial statements, the Company's compliance with legal or regulatory requirements related to the Company's financial statements,

reporting, internal control over financial reporting, or disclosure controls and procedures, the performance or independence of the Auditors, the performance of the Company's internal audit functions or such other matters as the Committee deems appropriate from time to time or whenever it shall be called upon to do so.

- (s) Review, discuss and assess at least annually its own performance as well as the Committee's role and responsibilities as outlined in this Charter.
- (t) Perform any other activities consistent with this Charter, the Company's Bylaws, and governing law as the Committee or the Board deems necessary or appropriate. Perform such other functions and have such powers as may be necessary or appropriate in the efficient and lawful discharge of the foregoing.

D. Breaches of Duties

51. Defendants' conduct constitutes a knowing and culpable breach of their duties as officers and/or directors of NVIDIA, demonstrating a lack of good faith and a reckless disregard for their obligations to the Company, as their actions violated privacy law and they either knew or were reckless in failing to recognize that their actions posed a significant risk of serious harm to the Company.

52. During the Relevant Period,² Defendants breached their duty of loyalty and

² The Relevant Period is January 2023 to the present.

good faith by causing the Company to adopt an unlawful business plan, pursuant to which they caused the Company to violate copyright laws, federal securities laws, BIPA, and to cause, or by themselves causing, the Company to make improper statements to the public and failing to implement adequate internal controls over legal and regulatory compliance. These improper practices caused NVIDIA to incur substantial damages.

53. Defendants, because of their positions of control and authority as officers or directors of NVIDIA, were able to and did, directly or indirectly, exercise control over the wrongful acts complained of herein. Defendants also failed to prevent the other Defendants from taking such illegal actions. As a result, and in addition to the damage the Company has already incurred, NVIDIA has expended, and will continue to expend, significant sums of money to remedy Defendants' failures.

V. FACTUAL BACKGROUND

A. The NVIDIA AI Process

54. NVIDIA's AI includes multiple generative-AI models and related tools and technologies. To train these generative-AI models, NVIDIA amassed an enormous library of data that includes copyrighted works, copied without author consent, credit, or compensation. NVIDIA's AI models include NVIDIA's LLMs, VLMs, and CVMs.

55. NVIDIA's LLMs are software programs trained by copying an enormous quantity of textual works and then feeding these copies into the model. These LLMs were trained on multiple pirated datasets and copied copyrighted works multiple times to train NVIDIA's language models, and were placed on the Hugging Face website.

56. NVIDIA's VLMs are large-scale generative text-to-video AI models that infringed copyrighted materials through pirated datasets comprised of YouTube videos and placed on Cosmos.

57. NVIDIA's CVMs are voice synthesis, voice cloning, and voice-conversion models that ingested hundreds of thousands of hours of human speech recordings, and extracted out source speakers' biometric signatures and voiceprints, in violation of BIPA. These speaker recording AI models were distributed through NVIDIA AI Enterprise, the NVIDIA API Catalog, NVIDIA NiM microservices, and open-weight model releases on Hugging Face.

B. NVIDIA's LLMs

58. In November 2021, NVIDIA announced its "NeMO Megatron framework for training language models," which NVIDIA touted as its "production-ready, enterprise-grade solution to simplify the development and deployment of" its LLMs.³ In September 2022, NVIDIA released its LLMs, NeMo Megatron-GPT 1.3B, NeMo Megatron-GPT 5B, NeMo Megatron-GPT 20B, and NeMo Megatron-T5 3B. Each of these models were hosted on the Hugging Face website.

59. On Hugging Face, NVIDIA's LLMs are described in "model cards" that provide information about each model. The model card for each of the NVIDIA LLMs states that "[t]he model was trained on 'The Pile' dataset prepared by EleutherAI." "The Pile" is a broader, publicly available dataset which contains Books3.

60. In October 2020, the Books3 dataset, consisting of approximately 196,640 books sourced from the online pirate "shadow library" Bibliotek, was made available online to provide AI developers with access to high-quality data. The Books3 dataset included copyrighted books without permission, in violation of the Copyright Act.

61. The Pile that includes Books3 is a dataset articulated in a paper by

³ Available at: <https://nvdianews.nvidia.com/news/nvidia-brings-large-language-ai-models-to-enterprise-worldwide>.

EluetherAI called “The Pile: An 800GB Dataset of Diverse Text for Language Modeling,” which described the Books3 as “a dataset of books derived from a copy of the contents of the Bibliotik private tracker ... Bibliotik consists of a mix of fiction and nonfiction books and is almost an order of magnitude larger than our next largest book dataset (BookCorpus2). We included Bibliotik because books are invaluable for long-range context modeling research and coherent storytelling.”

62. Bibliotek is well known as a notorious pirate website that contains vast quantities of unauthorized copyright materials used to train AI language models. Utilizing these illegal libraries provided NVIDIA and other AI entities with access to large quantities of training data rapidly. Such pirate libraries are in violation of the Copyright Act, and the AI developers’ use of such materials to train AI models is not shielded by the fair use doctrine.

63. The Books3 dataset, as confirmed by the person that assembled it, Shawn Presser, represents all of Bibliotek. And, downloading the Pile dataset from sources such as “The Eye” and the “Hugging Face” downloads a copy of Books3 and all of its unlicensed copyrighted material.

64. In October 2023, the Books3 dataset was removed based on the fact that it “is defunct and no longer accessible due to reported copyright infringement.” Books3 was included entirely with the RedPajama dataset, which is also available for downloading from the Hugging Face website. The Books3 dataset removal was messaged in RedPajama dataset documentation due to reported copyright infringement. Prior to October 2023, any person or entity that downloaded the RedPajama or The Pile datasets from Hugging Face was therefore downloading a copy of the Books3 dataset.

65. On March 20, 2025, the Danish Rights Alliance issued a report that

articulated “‘Classic pirate sources’ are widely used to train AI datasets.”⁴ “Generative AI model providers across the spectrum of LLMs, image, video and music generation (platforms) have obtained infringing copies of copyright protected works that were sourced from ‘classic pirate sources’ such as illegal filesharing and streaming sites.”⁵ The report also describes how NVIDIA, and other generative AI developers, obtain piracy datasets from Books3 and Common Crawl. Specifically, NVIDIA’s Nemo Megatron has used multiple publicly available complied datasets “that contain infringing copies of copyright protected works from illegal sources.”⁶

66. The Danish Rights Alliance Report also articulated that the Books3, Common Crawl, Opensubtitles, and RedPajama datasets with copyright infringement were contained in the Pile.⁷ This Report stated that NVIDIA’s Nemo Megatron used Books3, Opensubtitles, and Common Crawl, and “NVIDIA claims it has used The Pile (contains Books3, OpenSubtitles, Common Crawl) to train its Nemo Megatron-GPT 1.3B model.”

67. NVIDIA also downloaded the SlimPajama dataset that also included Books3. The SlimPajama dataset was created by copying and manipulating the RedPajama dataset (including copying Books3). NVIDIA used the SlimPajama dataset to test “both sentencepiece and BPE [tokenizers],” software used to process training data for use in NVIDIA’s LLM training and development. Stated by a NVIDIA employee: “SlimPajama ... is available in our org.”

⁴ Available at: <https://piracymonitor.org/report-classic-pirate-sources-are-widely-used-to-train-ai-datasets-says-danish-rights-alliance/>.

⁵ *Id.*

⁶ *Id.*

⁷ *Id.*

68. The SlimPajama dataset is also based in part on the Common Crawl dataset, which is known to contain copyrighted materials, including articles published on news websites and in digital newspapers. Also in 2023, a Dutch non-profit entity asked the creators of Common Crawl to remove over “two million news articles belonging to popular Dutch news outlets from its AI training dataset,” which were copied without permission.

69. The SlimPajama dataset contains copyrighted material that has not been licensed and authorized by its owners. Thus, NVIDIA’s usage of the SlimPajama dataset infringed copyright holders by downloading unauthorized copies of their works.

70. NVIDIA also coordinated with Anna’s Archive to access pirated copyrighted books on Anna’s Archive’s shadow libraries. These shadow libraries included Library Genesis (“LibGen”), Z-Library, and Sci-hub, all of which have been subject to assessment for copyright infringement by federal courts and/or law enforcement. Based on these shadow libraries, Anna’s Archive hosts millions of pirated books.

71. Anna’s Archive has disclosed “[i]t is well understood that LLMs thrive on high-quality data. We have the largest collection of books, papers, magazines, etc. in the world, which are some of the highest quality text sources.”⁸ Shadow libraries provide “high-speed ... enterprise-level access [to their collections] ... [in exchange] for donations in the range of tens of thousands USD.” These are payments for piracy.

72. Anna’s Archive also noted that the explosion in piracy and patronage by LLMs has saved shadow libraries from extinction:

Not too long ago, “shadow-libraries” were dying. Sci-Hub, the massive illegal archive of academic papers, had stopped taking in new works, due to lawsuits. “Z-Library”, the largest illegal library of books, saw its alleged creators arrested on criminal copyright charges ***Then came AI. Virtually all major companies building LLMs contacted us to train on our***

⁸ Available at: <https://annas-archive.org/llm>.

*data ... We have given high-speed access to about 30 companies.*⁹

73. NVIDIA reached out to Anna’s Archive in the fall of 2023 when it faced a rapidly approaching deadline in the form of its annual developer day. After OpenAI had released ChatGPT to massive success, NVIDIA sought to develop and demonstrate cutting edge LLMs for its NextLargeLLM. NVIDIA was “hyper [f]ocused on books corpuses,” and knew that “published books under copyright” are “the most valuable” for developing LLMs. NVIDIA contacted Anna’s Archive to discuss acquiring millions of pirated materials, “including Anna’s Archive in pre-training for our LLMs,” and sought to determine what “high-speed access” to the data would look like. Anna’s Archive informed NVIDIA that, because its collections were illegally acquired and maintained, NVIDIA executives would need to “let [Anna’s Archive] know when you have decided internally that this is something you can pursue. We have wasted too much time on people who could not get internal buy-in.”

74. Despite the warning from Anna’s Archive of the illegal nature of its collections, NVIDIA’s management gave “the green light” to proceed with piracy. Anna’s Archive offered millions of pirated copyrighted books, and access to several million books from Internet Archive, totaling roughly 500 terabytes of data. By downloading Anna’s Archive, NVIDIA pirated additional copies of copyrighted works, placed on the LLMs Nemotron-4 15B and Nemotron-4 340B in 2024.

C. NVIDIA’s VLMs

75. Cosmos is not a research tool or an isolated project. It is NVIDIA’s foundational video world model that provides a common layer of training data and video

⁹ Available at: <https://annas-archive.org/blog/ai-copyright.html>.

understanding to other services that include GR00T, Avatar, OV, and GeForce.¹⁰ NVIDIA has sought to develop Cosmos so it is capable, among other things, of turning language prompts into audiovisual content, and to create an AI product that accepts language instruction and produces a video based on that instruction. In order to adequately train its generative AI models, NVIDIA requires significant amounts of data to feed, train, improve, and commercialize Cosmos into a text-to-video generative AI model.

76. To access these datasets, NVIDIA's employees communicated via chats on the Slack platform a channel titled "#cosmos-dataset-creation," set up to "generate huge curated video dataset for video generative modeling." NVIDIA obtained datasets from a variety of sources. These datasets were treated by the NVIDIA employee group as "raw materials" for Cosmos even though most of the datasets contained licensed material for academic or non-commercial use that would involve downloading underlying copyrighted works.

77. NVIDIA was primarily focused on the use of HD-VG-130M, HDVILA100M, and HowTo100M datasets. Each of these datasets were all comprised of YouTube video source materials that must be downloaded from YouTube to be used for AI training to create the Cosmos foundational AI model. These datasets do not contain complete videos, instead specific "timestamps" within the videos divided up and designated as "clips." Each clip is treated as a separate training example and utilized as repeated retrieval for the same underlying video for each clip, which results in the same copyrighted work being copied multiple times. Because YouTube provides no lawful mechanism for downloading these works, every retrieval requires the user to circumvent

¹⁰ Samantha Cole, *Leaked Documents Show NVIDIA Scraping 'A Human Lifetime' of Videos Per Day to Train AI*, 404 MEDIA, Aug. 5, 2024 ("*NVIDIA Scraping*").

YouTube's technological restrictions, terms of service, and licensing limits. All downloading, copying, and circumvention occur on a clip-by-clip basis, and each instance constitutes its own violation.

78. Each of these datasets were comprised entirely of YouTube videos. HD-VG-130M contains 1,549,408 YouTube videos broken into 130 million high-definition video clips. HD-VILA-100M contains 3,098,462 YouTube videos divided into 100 million video clips. HowTo100M contains 1,238,911 YouTube videos split into 100 million short clips.

79. The HD-VG-130M dataset was technically created for academic research in 2023 and never set up for commercial use. Based on a platform set up by Microsoft, <https://github.com/daooshee/HD-VG-130M>, NVIDIA was aware that this dataset is comprised of protected, copyrighted works that can only be obtained without license, permission, or authorization. Despite this well-certified set of facts, NVIDIA used HD-VG-130M to retrieve every clip and incorporate those files into its commercial training platform. NVIDIA was aware that the datasets were for research purposes only but used the datasets for commercial development of its generative AI model, and to scrape audiovisual content from YouTube without the consent of YouTube or the creators of visual content.

80. The HD-VILA-100M dataset was compiled by Microsoft Research Asia and published in 2021. This dataset includes subtitled information extracted from YouTube's closed captioning system, and does not provide its own video files, so it must download each referenced clip directly from YouTube. Microsoft implicitly acknowledged that the HD-VILA-100M dataset is comprised of protected, copyrighted works obtained without license, permission, or authorization.

81. In addition, the HD-VILA-100M dataset is also compiled on the Panda-70M dataset, which utilizes 3.8M high-resolution long videos collected from HD-VILA-100M. The Panda-70M is a dataset compiled from 3,098,462 YouTube videos and consists of 70.8 million video clips extracted from those YouTube videos, none of which were obtained with the authorization of YouTube or the copyright holders of the YouTube videos. Panda 70M disclosed its dataset process in a paper entitled *Panda-70M: Captioning 70M Videos with Multiple Cross-Modality Teachers*:

We dub the dataset as Panda-70M. We show the value of the proposed dataset on three downstream tasks: video captioning, video and text retrieval, and text-driven video generation In this work, we present a large-scale dataset containing 70M video clips with caption annotations.

To establish the dataset with this mindset, we begin by using 3.8M high-resolution long videos collected from HD-VILA-100M and process them through the following three steps. First, we design a semantics-aware video splitting algorithm to cut long videos into semantically consistent clips while striking the balance between semantics coherence and the duration of the video clips. Second, we use a range of cross-modality teacher models, including image captioning models and image/video visual-question-answering (VQA) models with additional text inputs, such as video description and subtitles, to predict several candidate captions for a clip. Lastly, we collect a 100K video subset, where human annotators act as an oracle to select the best caption for each video. We use this dataset to finetune a fine-grained video-to-text retrieval model which is then applied to the whole dataset to select the most precise caption as the annotation.

82. In establishing its “video data curation pipeline” for the Cosmos World Foundation Model (“WFM”), and its Cosmos Tokenizer, NVIDIA “resort[s]” to existing video datasets that include Panda-70M.¹¹ “For Panda-70M, we manually filter out the videos” which results in a total of 500 videos for its “TokenBench.” Panda-70M contains

¹¹ NVIDIA, *Cosmos World Foundation Model Platform for Physical AI*, at 16, July 11, 2025; available at: <https://arxiv.org/pdf/2501.03575>.

models trained on 48 NVIDIA A100 GPUs.

83. NVIDIA was aware that the HD-VILA-100M dataset was for research purposes only, but nonetheless used it for commercial development of its AI model and to scrape audiovisual content from YouTube without the consent of YouTube. The *404 Media* article discloses how NVIDIA’s employees discussed downloading video content from links provided in Microsoft’s HD-VILA-100M dataset: “Found a list of files to download with high priority. Turns out there are -2.3M raw videos missing in the HDVILA dataset we have!” NVIDIA employees then sent a link to a Google Drive document to say “Here are the missing youtube links Let’s put this in the downloading pipeline!”

84. The HowTo100M dataset is hosted by European research institutes that created it. It contains only identifiers and metadata, but not any audiovisual material. The use of HowTo100M in training requires users to download each clip directly from YouTube. The dataset is labeled as a “commercial license,” but that has no effect on the legal rights associated with the underlying YouTube videos. Each of those videos remain protected by YouTube’s terms of service. Thus, every act of downloading a HowTo100M clip from YouTube is a separate unauthorized copying event and a separate circumvention violation, regardless of any license attached to the dataset file. NVIDIA carried out these downloads as part of its Cosmos training pipeline, resulting in more than 100 million additional unauthorized copying events.

85. Content creators that upload content to YouTube authorize and instruct YouTube to provide protection through YouTube’s anti-circumvention software, which is a driving factor to encourage content creators to upload their video content to YouTube. Rather than seek permission or pay a fair price for the audiovisual content hosted on YouTube, Defendants caused NVIDIA to harvest content creators’ protected and

copyrighted videos for commercial use and compilation into datasets without consent or compensation to the content creators.

86. YouTube permits public viewing of audiovisual works on its platform, but does not authorize access to audiovisual works through automated extraction, direct retrieval of media streams, or reconstruction of works outside its controlled environment. To prevent access to these video files, YouTube's TPMs are designed to control, restrict, and monitor access, and to deter extraction outside YouTube's controlled playback environment. YouTube's TPMs control how videos are delivered and accessed, providing videos only in small segment pieces, through temporary links placed only on YouTube's system. YouTube's TPMs then prevent users from obtaining the videos outside of YouTube's controlled delivery system.

87. To access this platform system, YouTube requires users to apply in the authorized process sequence, which includes "the issuance and validation of session-specific tokens, execution of player-side instructions, and reconstruction of segmented media streams into a continuous audiovisual experience." The audiovisual work cannot be accessed by public users without these technological processes.

88. YouTube's TPMs are "gatekeeping" mechanisms that include: "(1) an obfuscated signature system commonly referred to as a 'rolling cypher,' (2) IP-based blocking and rate limiting that restrict high-volume automated access, (3) short-lived, session-bound streaming URLs, (4) CAPTCHA human-verification challenges triggered by automated activity, and (5) proof-of-origin tokens that verify requests originating from authorized client environments."

89. NVIDIA's VLMs require significant amounts of data to feed, train, improve, and commercialize. To do so, Defendants caused NVIDIA to utilize automated

systems and tools that effectively bypassed YouTube's TPMs to access and obtain audiovisual works and reconstruct and assemble the protected media streams outside YouTube's authorized environment. NVIDIA then reproduced these videos to assemble training for NVIDIA's VLMs and services.

90. Defendants caused NVIDIA to obtain audiovisual datasets directly through its own independent scraping of YouTube's video sharing platform. NVIDIA's employees used tools and processes such as the open-source YouTube video downloader "yt-dlp." The yt-dlp downloader is a tool for downloading videos and audio from YouTube and other online platforms. NVIDIA used the yt-dlp downloader to improperly access and download YouTube files and merge data through the bypass of YouTube's player page, and avoid YouTube's monitoring systems, in order to scrape audio and video content from YouTube and feed content directly into its generative AI model.

91. YouTube allows users to "stream" content — playing it as it is retrieved — but it prohibits, through contractual and TPMs, making permanent, unrestricted copies of the content it streams. Bulk extraction of YouTube videos cannot occur without circumventing YouTube's TPMs, each of which independently controls access to the underlying audiovisual files. Illicit tools and services are specifically designed to defeat these protections, enabling automated systems to access content that YouTube makes available only for streaming. This process is commonly known as "scraping" or "stream ripping," and circumvention of any single technological measure that effectively controls access to a copyrighted work is sufficient to establish liability. YouTube's protections are layered and overlapping, and the circumvention of one does not negate the effectiveness or legal significance of the others.

92. NVIDIA also combined with virtual machines that refresh IP addresses.

This enabled NVIDIA to avoid being blocked by YouTube; refreshing IP addresses would *defeat* YouTube’s ability to: (i) monitor the downloading activity from IP addresses; and (ii) block those IP addresses from accessing audiovisual data on the YouTube platform.

93. Defendants did not require NVIDIA to obtain consent of YouTube to conduct its scraping activities. NVIDIA’s datasets with YouTube videos were scraped without permission from YouTube, and the scraping and acquisition processes used by Defendants to circumvent YouTube’s TPMs were inconsistent with, and in violation of, YouTube’s Terms of Service, which forbids scraping and mass downloading of videos.¹² NVIDIA’s automated system to access YouTube video content was intentionally designed and utilized to circumvent YouTube’s TPMs.

94. Defendants and NVIDIA’s primary employees were well aware that they had violated YouTube’s rules and the TPMs. As set forth in *NVIDIA Scraping*, NVIDIA “employees also discussed the issue of YouTube blocking IP addresses; if platforms detect something like a scraper being used to download mass amounts of content, they can block individual IP addresses from access.”¹³ One employee asked “re: YT blocking IPs have you considered something like <https://brightdata.com/> for IP rotation[]? We are considering this for scraping LLM data right now,” and another employee responded by stating “We are on AWS and restarting a VM instance gives new public IP. So that’s not a problem so far.”¹⁴

95. On the #cosmos-dataset-creation platform on Slack, NVIDIA employees

¹² Davey Alba & Emily Chang, *YouTube Says OpenAI Training Sora With Its Videos Would Break Rules*, BLOOMBERG, Apr. 5, 2024.

¹³ *NVIDIA Scraping*, at 10-11.

¹⁴ *Id.* at 11.

“occasionally brought up legal and ethical questions about what they’re working on.”¹⁵ In February 2024, NVIDIA’s principal scientist, Francesco Ferroni, asked if NVIDIA “cannot use [YouTube-8M] for non-research purposes [for the Cosmos project]?” An employee responded by stating that “downloading YouTube videos is prohibited by YouTube’s TOS,” but “for downloading YouTube 8m, we cleared the download with Google/YouTube ahead of time and dangled as a carrot that we were going to do so using Google Cloud.” Ferroni also asked if NVIDIA has “an API account with YouTube? We will need to check the licenses of videos, both from open-source datasets, as well as any other videos we source via that.” Ming-Yu Liu, Vice President of Research at NVIDIA and a Cosmos project leader responded: “Let’s create an account? I downloaded 8 million YouTube videos in the past (YouTube 8m). I cleared this with YouTube ahead of time. I used about 10000 cores and thousands of IP addresses in Google Cloud to do it. I don’t think you need to be logged in to download videos.”

96. Ferroni also asked “presumably this data can only be used for research purposes?” “Can you also comment on the usage terms for ... YouTube8M?” A NVIDIA employee responded by stating:

I don’t know what license terms Google filtered on when creating the dataset: we just downloaded whatever they listed as being contained in the dataset The YouTube 8m dat[a] that I downloaded was downloaded with complete metadata public domain is, of course, always best. However, whether using copyrighted material is fair use is currently an open legal issue [but] our legal team has OK’ed this kind of thing for LLM training and may OK it for video training as well.”¹⁶

97. Another employee asked: “Hi team. Are we using

¹⁵ *Id.*

¹⁶ *Id.* at 12-13.

<https://research.google.com/youtube8m/download.html> to download the videos? If yes, do we have legal approval for the same? In one of the projects, legal denied using it because license on the individual videos supersede the license shared on yt8m.” Liu responded by stating “This is an executive decision. We have umbrella approval of all the data.”¹⁷

98. Accordingly, Defendants were aware of the legality issues of allowing NVIDIA to use the YouTube video dataset to train its generative AI models. Regardless of their realization that they had permitted NVIDIA to engage in potential copyright infringement, Defendants allowed NVIDIA to continue its program of improperly accessing YouTube data and using the audiovisual content NVIDIA improperly accessed to train its generative AI models. Based on this illegality concession, in March 2024 the Cosmos “project hit a milestone: 100,000 videos downloaded, which NVIDIA accomplished in 2 weeks.” “In late May [2024], an email about data strategy ... announc[ed] that [NVIDIA] compiled 38.5 million video URLs.”¹⁸

D. NVIDIA’s CVMs

99. NVIDIA built a suite of commercial voice AI products using the voices of real people, establishing its CVMs, including Magpie TTS Zeroshot, Magpie TTS Flow, FUGATTO, PersonaPlex, Nemotron 4 VoiceChat, and the Canary and Parakeet speech models. NVIDIA’s voice AI research is conducted primarily through its Applied Deep Learning Research group (“ADLR”). Prior to the CVMs, ADLR and NVIDIA research groups released a sustained line of voice synthesis, voice cloning, and voice-conversion models: WaveGlow (2018), Mellotron (2019), Flowtron (2020), FastPitch (2020), RAD-TTS (2021-22), Mixer-TTS (2022), and Incremental FastPitch (2023-24). NVIDIA

¹⁷ *Id.* at 14-15.

¹⁸ *Id.* at 18-19.

distributes its CVMs through NVIDIA AI Enterprise, the NVIDIA API Catalog, NVIDIA NiM microservices, and open-weight models releases on Hugging Face. NVIDIA's foundational voice synthesis model occurred on NVIDIA's operated computing infrastructure, including on NVIDIA's GPUs.

100. NVIDIA has disclosed how the AI voice synthesis process works by training neural networks on large quantities of recorded human speech. The networks identify and reproduce acoustic features that distinguish individual voices: pitch, timber, resonance, accent, cadence, articulation, and the dynamics of emotional expression. These features are encoded on the network and stored on the model's parameters and then used to generate new speech that exhibits the acoustic features of the training voices.

101. In 2018, NVIDIA published the WaveGlow paper that introduced a flow-based generative network for raw audio synthesis trained directly on speech recordings.¹⁹

The WaveGlow paper stated:

As voice interactions with machines become increasingly useful, efficiently synthesizing high quality speech becomes increasingly important Speech synthesis requires generating very high dimensional samples with strong long term dependencies [R]eal time speech has challenging speed and computational constraints and higher sampling rates generate even higher quality speech.

[S]tate of the art speech synthesis models are based on parametric neural networks. Text-to-speech synthesis is typically done in two steps. The first step transforms the text into time-aligned features [the] second model transforms these time-aligned features into audio samples.

¹⁹ Ryan Prenger, Rafael Valle & Bryan Catanzaro, *WaveGlow: A Flow-based Generative Network for Speech Synthesis*, ARXIV: 1811.0002 (Oct. 31, 2018), available at: <https://arxiv.org/abs/1811.00002>.

Our contribution is a flow-based network capable of generating high quality speech We refer to this network as WaveGlow, as it combines ideas from Glow and WaveNet. WaveGlow is simple to implement and train, using only a single network, trained using only the likelihood loss function Mean Opinion Scores show that it delivers audio quality as good as the best publicly available WaveNet implementation trained on the same dataset.

WaveGlow is a generative model that generates audio by sampling from a distribution. To use a neural network as a generative model, we take samples from a simple distribution and put those samples through a series of layers that transform the simple distribution to one which has the desired distribution.

For all the experiments we trained on the LJ speech data. This data set consists of 13,100 short audio clips of a single speaker reading passages from 7 non-fiction books. The data consists of roughly 24 hours of speech data recorded on a MacBook Pro using its built-in microphone in a home environment. We use a sampling rate of 22,050kHz.

The WaveGlow network was trained on 8 Nvidia GV100 GPU's using randomly chosen clips of 16,000 samples for 580,000 iterations.

Existing neural network based approaches to speech synthesis fall into two groups. The first group conditions future audio samples on previous samples in order to model long term dependencies. The first of these auto-regressive neural network models was WaveNet which produce high quality audio.

WaveGlow networks enable efficient speech synthesis with a simple model that is easy to train. We believe that this will help in the deployment of high quality audio synthesis.

102. In 2021-22, NVIDIA's RAD-TTS research group disclosed the disentanglement of voice characteristics from linguistic content on the same model that can generate the same words in different speakers' voices, and the same speakers saying different words. In the "RAD-TTS: Parallel Flow-Based TTS with Robust Alignment

Learning and Diverse Synthesis” paper published in 2021, NVIDIA employees stated that the RAD-TTS model “extends prior parallel approaches by additionally modeling speech rhythm as a separate generative distribution to facilitate variable token duration during inference [and] feature[s] robust alignment learning and diverse synthesis.” NVIDIA’s “work extends their alignment setup to be more generally applicable to arbitrary TTS frameworks, as well as with improved stability. We further tackle the limited synthesis diversity in parallel TTS architectures. These architectures first determine the durations of each phoneme. Then, a parallel architecture maps the phonemes, replicated time-wise based on their predicted durations.” NVIDIA “aims to construct a generative model for sampling mel-spectrograms given text and speaker information ... [and] devise[s] an unsupervised learning approach that learns the alignment between text and speech rapidly without depending on external aligners. The alignment converges rapidly towards a usable state in a few thousand iterations.”

103. In the “Multilingual Multiaccented Multispeaker TTS With RADTTS” paper, NVIDIA employees disclosed that NVIDIA:

[W]orks to create a multilingual speech synthesis system which can generate speech with the proper accent while retaining the characteristics of an individual voice. This is challenging to do because it is expensive to obtain bilingual training data in multiple languages, and the lack of such data results in strong correlations that entangle speaker, language, and accent, resulting in poor transfer capabilities. To overcome this, we present a multilingual, multiaccented, multispeaker speech synthesis model based on RADTTS with explicit control over accent, language, speaker and fine-grained F0 and energy features. Our proposed model does not rely on bilingual training data. We demonstrate an ability to control synthesized accent for any speaker in an open-source dataset comprising of 7 accents. Human subjective evaluation demonstrates that our model can better retain a speaker’s voice and accent quality than controlled baselines while synthesizing fluent speech in all target languages and accents in our dataset.

In this work, we (1) demonstrate effective scaling of single language TTS to multiple languages using a shared alphabet set and alignment learning framework; (2) introduce explicit accent conditioning to control the synthesized accent; (3) propose and analyze several strategies to disentangle attributes (speaker, accent, language and text) without relying on parallel training data (multilingual speakers); and (4) explore fine-grained control of speech attributes such as F0 and energy and its effects on speaker timbre retention and accent quality.

We present a multilingual, multiaccented and multispeaker TTS model based on RADTTS with novel modifications. We propose and explore several disentanglement strategies resulting in a model that improves speaker, accent and text disentanglement, allowing for synthesis of a speaker with closer to native fluency in a desired language without multilingual speakers. Internal ablation studies indicate that explicitly conditioning on fine-grained features (F0 and E) results in better speaker retention and pronunciation according to human evaluators. Our model provides an ability to predict such fine-grained features for any desired combination of speaker, accent and language and user studies show that under limited data constraints, it improves pronunciation in novel languages. Scaling the model to large-resource conditions with more speakers per accent remains the subject of future work.

104. In October 2025, NVIDIA employees described and disclosed the Magpie TTS models, Magpie TTS Multilingual, Magpie TTS Zeroshot, and Magpie TTS Flow, in the blog post “Enhancing Multilingual Human-Like Speech and Voice Cloning with NVIDIA Riva TTS”:

NVIDIA Riva is a suite of multilingual microservices for building real-time speech AI pipelines. Riva delivers top-notch accuracy in TTS, ASR, and neural machine translation (NMT), and works across on-prem, cloud, edge, and embedded devices. TTS, also called speech synthesis, converts text into high-quality, natural-sounding speech. It has been a challenging task in the field of speech AI for decades.

The Magpie TTS Multilingual and Magpie TTS Zeroshot models build on an encoder-decoder transformer architecture to target streaming applications The output of the model is the generated acoustic tokens of the target speaker. The two models are built on different variants of the encoder-decoder architecture:

Magpie TTS Multilingual: Uses a multi-encoder setup. A dedicated context encoder processes context audio tokens. The context and text encoder outputs are fed into different layers of the AR decoder. This allows for clear separation of modalities.

Magpie TTS Zeroshot: Employs a decoder context setup, utilizing the decoder's self-attention mechanism for speaker conditioning. It feeds context audio tokens into the AR decoder, which then processes both the context and target audio tokens. This approach uses a shared representation for both conditioning and prediction.

You can find the voice names and emotions the models currently support in the pretrained TTS models documentation.

The input for this pipeline has two parts: a user prompt, and a target speaker's audio prompt. An LLM hosted by NVIDIA LLM NIM generates random text for TTS tasks based on a user prompt. The Magpie TTS Zeroshot model takes the text and audio prompt as input, then generates the corresponding audio of the input text with the target speaker's voice.

The Magpie TTS Flow model introduces an alignment-aware pretraining framework that integrates discrete speech units (HuBERT) into an NAR training framework (E2 TTS) to learn text-speech alignment.

Magpie TTS Flow integrates alignment learning directly into the pretraining process using untranscribed speech data, without the need for separate alignment mechanisms. By embedding alignment learning into pretraining, it facilitates alignment-free voice conversion and allows for faster convergence during fine-tuning even with limited transcribed data.

Magpie TTS Flow can achieve high pronunciation accuracy (lower WER) and higher speaker similarity (SECS-O) with significantly fewer pretraining and fine-tuning iterations compared to other models. Furthermore, it can also learn text-speech alignment effectively for multiple languages by adding language ID as an input to the decoder, making it a robust multilingual TTS system the released Riva model was trained on a significantly larger paired dataset (about 70K hours) to further improve zero-shot performance. You can synthesize the voices of target speakers who have consented to such use Magpie TTS Flow takes the inputs, then

generates the corresponding audio of the input text with the target speaker's voice.

105. On the Hugging Face platform, NVIDIA's employees described NVIDIA AI dataset to be conditioned on a "voice prompt" that consisted of audio tokens establishing "target[ed] vocal characteristics and speaking style" with voice data to clone a speaker's voice resulting in a voice print. The voice print established BIPA liability in the context of a "biometric identifier," which includes audio tokens, speaker embeddings, voice prompts, and speaker-acoustic-signature representations produced by NVIDIA's voice synthesis pipeline. This biometric processing occurred in Magpie TTS Zeroshot, Magpie TTS Flow, PersonaPlex, and during base-model training.

106. NVIDIA announced FUGATTO in November 2024, an audio generative model trained in excess of 50,000 hours of audio, capable of transforming human voices and modifying accent and emotion. NVIDIA released Magpie TTS Zeroshot and Magpie TTS Flow in October 2025, and PersonaPlex in January 2026.

107. In complete opposition to NVIDIA's statements that a voice-fraud detection company, "Pindrop," would develop deepfake detection capabilities against NVIDIA's voice models, that voice cloning would only be used with "target speakers who consented to such use, and that NVIDIA's Trustworthy AI terms would prohibit use of NVIDIA's models for "unlawful biometric data collection," NVIDIA's foundational voice synthesis models were trained on voiceprints extracted without the knowledge or consent of the speakers whose voice recordings were ingested during training.

108. NVIDIA stated that it drew its datasets from public-domain sources, including LibriVox audiobooks, Project Gutenberg texts, and LibriTTS as an open, publicly available speech dataset. But there are no actual claims or evidence that the approximately 70,000 hours of paired speech data on Magpie TTS, or the 50,000 hours of

audio on FUGATTO, were trained entirely on properly licensed or consented data. Instead, NVIDIA obtained substantial portions of its training data, across multiple foundational-model families, from publicly accessible internet sources without the consent of source speakers, authors, or content creators.

109. In August 2025, NVIDIA released the Granary dataset, which included approximately a one-million-hour corpus of speech audio assembled by processing pre-existing open-source corpora (YODAS, YouTube-Commons, VoxPopuli, LibriLight) through NVIDIA's NeMo Speech Data processor pipeline. The Granary dataset was used to train NVIDIA's Canary-1b-v2 and Parakeet-tdt-0.6b-v3 speech models distributed on Hugging Face. These one-million-hours of speech audio processing were placed on platforms without the consent of speakers whose voices appeared in the recordings.

110. Based on this process, NVIDIA distributes these voice and speech synthesis models on Hugging Face. Once published on Hugging Face, these models are disseminated to third parties that NVIDIA cannot recall from the downloads, which places the voiceprints and biometric information of source speakers, which implicates their legal rights established by BIPA. In violation of BIPA, NVIDIA failed to identify the source speakers, provide written notice of the specific purpose and duration of collection, and obtain a written release from each speaker before ingesting that speaker's recording into the training pipeline, which resulted in damages to the source speakers.

E. Red Flags of NVIDIA's Copyright Infringement

111. Defendants were well aware of the potential liability stemming from NVIDIA's unauthorized use of copyrighted works in developing NVIDIA's LLMs, VLMs, and CVMs.

///

112. First, all the Officer Defendants were personally involved in: (i) the development of NVIDIA's AI strategy and the details of the NeMo Megatron models NVIDIA chose as its LLMs to develop the Company's AI products instead of other commercially available LLMs that did not include copyrighted material; (ii) the unauthorized access and extraction of videos from the YouTube platform to obtain the training data necessary to build and commercialize its NVIDIA AI Video, Cosmos, through the use of HD-VG-130M, HDVILA100M, and HowTo100M datasets; and (iii) the utilization of NVIDIA's CVMs, primarily Magpie TTS Zeroshot, Magpie TTS Flow, FUGATTO, PersonaPlex, Nemotron VoiceChat, Canary, and Parakeet, which ingested hundreds of thousands of human speech recordings placed on Hugging Face, and failed for NVIDIA to provide any notice to, or release from, any of the speakers whose recordings were placed on the training pipeline. The Director Defendants were made aware of the Officer Defendants' engagement in this copyright infringement process.

113. Second, the multiple individual and class action lawsuits against other AI developers for copyright infringement put all Defendants on notice of the need to avoid using unlicensed copyrighted materials in the LLMs, VLMs, and CVMs used to train NVIDIA Intelligence AI models. Multiple class actions filed beginning in 2023, prior to the copyright infringement lawsuits filed against NVIDIA, involve claims against NVIDIA's competitors based on the exact same approach to Defendants' construction of an AI training database premised on multiple "shadow libraries," including The Pile, Books3, Anna's Archive, RedPajama, LibGen, SciHub, Z-Library, and Panda-70M, datasets comprised of YouTube videos, and human speech recordings.

///

///

114. The first action filed against Anthropic in 2023, *Concord I*, alleged that Anthropic “torrented ‘countless pirated copies’ of songbooks and sheet music.”²⁰ Anthropic “used BitTorrent to download publishers’ works” and “simultaneously uploaded to the public large unauthorized copies of the same works, separately violating publishers’ exclusive right of distribution and encouraging further infringement of their copyrighted works.”²¹

115. In the *Mosaic* action, copyright holders filed a lawsuit against Mosaic ML and Databricks in 2024. *See In re Mosaic LLM Litigation*, No. 3:24-cv-01451-CRB (N.D. Cal. 2024). Those copyright holders allege damages from use of their works in a “Mosaic ML”. That dataset was used in developing the MosaicML Pretrained Transformer, a type of artificial intelligence/LLM software.

116. Defendants knew or recklessly disregarded this vortex of litigation targeting AI developers, all of which was widely reported. Defendants were also on notice of the widely reported resolution of the *Concord I* case, for which Anthropic paid damages of \$1.5 **billion** to resolve after Judge Alsup on June 23, 2025 issued a summary judgment Order on Fair Use establishing Anthropic’s potential liability for the “Pirated Library Copies.” The Court “denie[d] summary judgment for Anthropic that the pirated library copies must be treated as training copies,” and set up a “trial on the pirated copies used to create Anthropic’s central library and the resulting damages, actual or statutory (including for willfulness).” *Concord I*, No. 3:24-cv-05417-WHA (N.D. Cal.) (ECF No. 231 at 31-32).

²⁰ Hailey Konnath, *Anthropic Hit With 2nd Music IP Suit, This Time for \$3B*, LAW360 (Jan. 28, 2026); available at: <https://www.law360.com/articles/2435463/>.

²¹ *Id.*

117. All Defendants were well aware that the unlawful business model they caused NVIDIA to adopt involved NVIDIA's violation of copyright laws premised on copyrighted book holders long before *Nazemian* filed against NVIDIA. As early as the October 2023 reported "removal" of the Books3 dataset "due to reported copyright infringement," Defendants knew that NVIDIA was utilizing pirated library copies and would be subject to billions of dollars of damages. These copyrighted materials were placed on the RedPajama or the Pile datasets, and both of these datasets still maintain the copyrighted materials downloaded from the Books3 dataset.

118. In addition, multiple BIPA lawsuits were filed in accord with the BIPA lawsuit filed against Apple, including: *Amer v. Eleven Labs Inc.*, No. 1:26-cv-05437 (N.D. Ill. 2026); *Marin v. Alphabet, Inc.*, No. 1:26-cv-05436 (N.D. Ill. 2026); *Marin v. Meta Platforms, Inc.*, No. 1:26-cv-05438 (N.D. Ill. 2026); and *Flowers v. Microsoft Corp.*, No. 1:26-cv-05491 (N.D. Ill. 2026).

119. Indeed, another AI development company, Anthropic, paid damages of \$1.5 billion to resolve the *Concord I* action and is now subject to additional damages of up to \$3 billion in a second action filed in 2026, *Concord II*, No. 5:26-cv-00880-EKL (N.D. Cal. 2026). This tidal wave of litigation filed against AI developers like NVIDIA served as blazing red flags alerting Defendants of the need to prevent NVIDIA from committing and continuing to commit the same unlawful conduct.

120. Copyright holders and commercial voiceprint holders have filed three class actions against NVIDIA for copyright infringement and BIPA violations.

///

///

F. The Copyright Infringement and BIPA Class Actions Against NVIDIA

121. Copyright holders and commercial voiceprint holders have filed three primary class actions against NVIDIA for copyright infringement and BIPA violations.²² On March 8, 2024, copyrighted book authors filed their class action lawsuit against NVIDIA, *Nazemian v. NVIDIA Corp.*, No. 4:24-cv-01454-JST (N.D. Cal.). On November 26, 2025, intellectual property video creators filed their class action lawsuit against NVIDIA, *Ted Entertainment, Inc. v. NVIDIA Corp.*, No. 5:25-cv-10287-EJD (N.D. Cal.). On May 12, 2026, commercial voice print creators filed their class action lawsuit against NVIDIA, *Rogers v. NVIDIA Corp.*, No. 1:26-cv-05478 (N.D. Ill.).

122. The *Nazemian* action, initially filed in March 2024, was amended into a First Consolidated Amended Complaint on January 16, 2026 (the “*Nazemian* Complaint”) (ECF No. 235). The *Nazemian* Complaint alleged that NVIDIA “copied ... copyrighted works multiple times to train its language models, including from known pirated libraries (also known as “shadow libraries”). Those notorious shadow libraries include The Pile, Bibliotik, and Anna’s Archive.” ECF No. 235 at 2. “NVIDIA and its employees subsequently made additional unlawful copies of this illegally-obtained copyrighted

²² On January 29, 2026, an additional copyright infringement action was filed based on NVIDIA’s use of “YouTube-hosted audiovisual works ... to train, develop, and improve NVIDIA’s ... Cosmos.” *Youngblood v. NVIDIA Corp.*, No. 5:26-cv-00916-NC (N.D. Cal. 2026). On March 26, 2026, an additional copyright infringement action was filed based on NVIDIA’s use of Microsoft’s TRELLIS-500K dataset to train NVIDIA “commercial AI infrastructure” incorporated into “AI training datasets.” *Beaulier v. NVIDIA Corp.*, No. 5:26-cv-02647-EKL (N.D. Cal. 2026). On June 22, 2026, an additional copyright infringement action with additional breach of contract, unjust enrichment, and unfair business practice claims was filed “arising from [NVIDIA’s] unauthorized use of the Jamendo database as well as the musical works and sound recordings ... in the development, training, and operation of a commercial generative AI system.” *S.A. Jamendo v. NVIDIA Corp.*, No. 5:26-cv-06206-NW (N.D. Cal. 2026) (also filed in Belgium).

material during the LLM development process.” *Id.* at 5. Each of NVIDIA’s NeMo Megatron models “were hosted on a website call Hugging Face,” which disclosed that “[t]he model was trained on “The Pile” dataset prepared by EleutherAI.” *Id.* at 7. “NVIDIA also downloaded the SlimPajama dataset” and “made unlawful copies of The Pile” and “sought vastly more copyrighted works than The Pile could provide” utilizing additional shadow libraries” that included LibGen, Z-Library, Sci-hub, and Anna’s Archive. *Id.* at 9-10.

123. Further, the *Nazemian* Complaint disclosed, based on the discovery process in the action:

Internal documents show competitive pressures drove NVIDIA to piracy.

In the fall of 2023, NVIDIA faced a rapidly approaching deadline in the form of its annual developer day. In the year since the launch of the NeMo Megatron series in September 2022, OpenAI had released ChatGPT to massive success, resulting in a substantial increase in investor attention on AI. In response, NVIDIA sought to develop and demonstrate cutting edge LLMs at its fall 2023 developer day. In seeking to acquire data for what it internally called “NextLargeLLM,” “NextLLMLarge” and “Next Generation LLM” (collectively, “NextLargeLLM”). NVIDIA was “[h]yper [f]ocused on books corpuses.” *NVIDIA knew that “published books under copyright” are “the most valuable” for developing LLMs and NVIDIA knew that only books were available in sufficient quantities.* And NVIDIA needed to achieve 8 trillion tokens for the “NextLargeLLM,” and books provided this means.

In August 2023, NVIDIA contacted books publishers to obtain fast “access to large volumes of unique, high-quality datasets” or “i.e. books.” But on information and belief, NVIDIA could not secure this fast access to the huge quantity of books it needed through publishers. As one book publisher told NVIDIA, it was “not in a position to engage directly just yet but will be in touch.” *In 2023, NVIDIA had “chatted with multiple publishers . . . but none [] wanted to enter into data licensing deals.”*

Desperate for books, NVIDIA contacted Anna’s Archive—the largest and most brazen of the remaining shadow libraries—about acquiring its millions of pirated materials and “including Anna’s Archive in pre-training data for our LLMs.”

Within a week of contacting Anna’s Archive, and days after being warned by Anna’s Archive of the illegal nature of their collections, NVIDIA management gave “the green light” to proceed with the piracy. Anna’s Archive offered NVIDIA millions of pirated copyrighted books. Anna’s Archive also offered access to several million books from Internet Archive, which were only normally available through Internet Archive’s digital lending system (a system which was found to be copyright infringement by the Second Circuit, *see Hachette Book Grp., Inc. v. Internet Archive*, 115 F.4th 163 (2d Cir. 2024)). Anna’s Archive promised NVIDIA access to “a lot of books,” totaling roughly 500 terabytes of data. By downloading Anna’s Archive, NVIDIA pirated additional copies of Plaintiff’s Infringed Works.²³

124. On January 29, 2026, NVIDIA, the defendant in the *Nazemian* action, moved to dismiss the *Nazemian* Complaint. The court essentially denied the motion on May 5, 2026.²⁴ The court denied the motion to dismiss based on: (i) “Megatron 354M” copyright infringement; (ii) “any unidentified models”; (iii) NVIDIA’s use of datasets including Bibliotek and BitTorrent Protocol; (iv) third party infringement; and (v) contributory infringement through “NVIDIA’s assistance to specific customers who used NVIDIA scripts to download The Pile” via NeMo.

125. As described in the *Law360* news report by Elliot Weld, “NVIDIA Must Face Most of Authors’ AI Copyright Suit,” dated May 6, 2026, NVIDIA was unsuccessful in its attempt to obtain dismissal of the copyright infringement claims:

U.S. District Judge Jon Tigar in a Tuesday order partly granted and partly denied Nvidia's motion to dismiss claims from authors Abdi Nazemian, Brian Keene and Stewart O’Nan. The writers sued Nvidia in March 2024, alleging the company trained multiple large language models in the Megatron family on materials taken from shadow libraries, online repositories that give consumers open access to works they would otherwise need to pay or register for.

²³ Emphasis added.

²⁴ Order Granting in Part and Denying in Part Defendants’ Motion to Dismiss, *Nazemian*, 4:24-cv-01454-JST (N.D. Cal. May 5, 2026).

Specifically, Nvidia asked Judge Tigar to toss allegations related to its Megatron 345M model and dismiss references to other “unidentified models” from the suit. The authors alleged Megatron 345M was trained on a dataset called “The Pile,” which contains a subset of pirated books called “Books3,” including the plaintiffs’ works. Judge Tigar said Books3 constitutes about 12% of The Pile, and at this stage, the plaintiffs had adequately alleged that Megatron 345M was trained on a dataset containing their works.

“Courts have routinely declined to dismiss allegations where plaintiffs alleged that their copyrighted material was present in datasets that contained their work,” Judge Tigar wrote.

On the issue of “unidentified models,” the judge said it became clear at a hearing on the motion to dismiss that the authors are not presently accusing any unidentified model of infringement.

Nvidia also asked to dismiss allegations regarding shadow libraries Pirate Library Mirror and Bibliotik, but the judge denied this part of the motion and said references to those two in the suit were largely historical and not directed at Nvidia.

Nvidia also wanted to dismiss allegations that it used the BitTorrent Protocol, a peer-to-peer file sharing tool. Judge Tigar declined to do so as BitTorrent is merely a means of downloading, saying “asking to dismiss allegations concerning BitTorrent is like asking to dismiss allegations concerning paintbrushes in a case about a dolphin painting.”

On the claims of inducing direct infringement by third parties, which stem from allegations that Nvidia provided customers with scripts to automatically download The Pile, Judge Tigar found the authors had sufficiently alleged predicate acts of direct infringement.

126. The *Ted Entertainment* action, filed on November 26, 2025, alleged that NVIDIA “access[ed] and scrape[d] millions of copyrighted videos from ... YouTube, in order to feed, train, improve and commercialize” NVIDIA’s AI “model named ‘Cosmos.’” ECF No. 1 at 1. *Ted Entertainment* articulated:

By scraping and downloading those files to build Cosmos, Defendant deliberately circumvented YouTube’s access controls to create a core AI infrastructure that would power Defendant’s entire ecosystem.

Rather than seek permission or pay a fair price for the audiovisual content hosted on YouTube, Defendant harvested content creators' protected and copyrighted videos for commercial use and at scale without consent or compensation.

Defendant's actions were not only unlawful, but an unconscionable attack on the community of content creators whose content is used to fuel the multi-trillion-dollar generative AI industry without any compensation.

Content creators such as Plaintiffs and the Class Members will never be able to clawback the intellectual property unlawfully copied and used by Defendant to train its generative AI. Once AI ingests content, that content is stored in its neural network and not capable of deletion or retraction. Defendant's actions constitute abuse and exploitation of content creators' work for Defendant's profit.

Cosmos was not a research tool or isolated project; it was Defendant's foundational video world model. Defendant's own internal materials show that Cosmos was designed to sit at the center of Defendant's architecture, providing a common layer of training data and video understanding that other services—such as GR00T, Avatar, OV, and GeForce—would rely upon[.]

Because Cosmos was foundational, Defendant had an overwhelming incentive to acquire training data on an unprecedented scale. Rather than negotiate for lawful licenses, Defendant broke through YouTube's access protections to obtain the massive datasets necessary to fuel Cosmos and, by extension, Defendant's broader product lines.

Defendant obtained datasets from a variety of sources, including academic repositories, research compilations, and other large-scale video collections created by universities, corporations, and independent researchers. These datasets were treated by Defendant as raw material for Cosmos, even when the datasets were expressly licensed for academic or non-commercial use and prohibited commercial exploitation, redistribution, or any use that would involve downloading the underlying copyrighted works.

This case is about Defendant's use of the HD-VG-130M, HDVILA100M, and HowTo100M datasets. Each comprised of YouTube videos ingested to create the Cosmos foundational AI model.

When Defendant scraped audio and video files from YouTube, Defendant did not simply download those files onto the YouTube app for offline streaming, as envisioned by YouTube's Premium plan. Instead, Defendant improperly accessed the actual audio and video files and downloaded those files into Defendant's own system, where Defendant had control over the files and where Defendant could store them indefinitely. These actions are inconsistent with Defendant's Premium plan.

127. The *Ted Entertainment* action filed a First Amended Class Action Complaint on March 16, 2026, in response to a motion to dismiss filed by NVIDIA on February 23, 2026. The *Ted Entertainment* amended complaint directly addressed and undermined the arguments made by NVIDIA in its motion to dismiss, essentially establishing that NVIDIA circumvented the YouTube video platforms to access controls and not just to copy controls, which confirms with liability pursuant to the Digital Millenium Copyright Act, 17 U.S.C. § 1201(a).

128. The *Rogers* action, filed on May 12, 2026, alleged that "NVIDIA built a suite of commercial voices of real people" to "power Magpie TTS Multilingual, Magpie TTS Zeroshot, Magpie TTS Flow, FUGATTO, PersonalPlex, Nemotron 3 VoiceChat, and the Chancery and Parakeet speech models ... to power ... NVIDIA AI Enterprise, the NVIDIA API Catalog, NVIDIA NIM microservices, and open-weight model releases on Hugging Face — NVIDIA ingested hundreds of thousands of hours of human speech and extracted the unique biometric signatures, the voiceprints, of the speakers from those recordings." NVIDIA failed to comply with BIPA, which:

[W]ould have required NVIDIA to identify the source speakers, provide written notice of the specific purpose and duration of collection, and obtain a written release from each speaker before ingesting that speaker's recording into the training pipeline. With foundational models trained on a corpus measured in hundreds of thousands of hours of speech and millions of distinct speakers, that compliance burden would have constrained the speed and scale of NVIDIA's voice AI development. NVIDIA chose speed and scale over compliance.

The voiceprints NVIDIA extracted from Plaintiffs are not stored in a database that can be deleted on request. They are encoded in the parameters of NVIDIA’s commercial voice models and reproduced in the audio those models generate. NVIDIA has also published the parameters of several of those models — including PersonaPlex, Canary, and Parakeet — as open-weight releases on Hugging Face under the NVIDIA AI Open Model License and the NVIDIA Nemotron Open Model License, available for download by any party worldwide. At this point, the biometric data and the product are the same thing. The product has already been distributed.

G. Multiple Copyright Infringement Actions Filed in Accord with Defendants’ Allowance of Copyright Infringement

129. In addition to the *Nazemian* action filed in March 2024, copyright holders have pursued multiple litigations against AI developers focused on the copyright-infringing data set up on the pirated shadow libraries, including RedPajama, Books3, The Pile, Bibliotek, and Anna’s Archive.

130. Copyright holders filed a lawsuit against Snowflake in 2025, alleging that RedPajama is a training database that “‘contained within it a deduplicated copy’ of Books3... ‘[a] dataset described in [the Pile paper as] derived from a copy the of contents of [] Bibliotek.’” *James v. Snowflake Inc.*, No. 2:25-cv-00108-BMM (D. Mont. 2025) (ECF No. 1 at 4). Copyright holders there alleged that to “train [its] Arctic [models], Snowflake downloaded, copied, stored, and used the RedPajama dataset [that] contain[ed] Books3... and repeatedly ... processed [] works during preprocessing and pretraining, and retained copies ... on its servers for further training.” *Id.* at 6.

131. Copyright holders also filed an action against Salesforce in 2025 for improperly training its artificial intelligence models on copyrighted works and relied on “notorious” datasets of pirated books to create its AI models. *Tanzer v. Salesforce Inc.*, No. 3:25-cv-08862-CB (N.D. Cal. 2025) ECF No. 1 at 1. Salesforce also used the same two datasets that NVIDIA used, the RedPajama and the Pile, which “contain hundreds of thousands of copyrighted books that were acquired without the authorization or consent of

the authors.” *Id.*

132. Copyright holders filed an action against Microsoft in 2025 for copying “a notorious collection of approximately 200,000 pirated books known as ‘Books3’ and fed them its [Megatron] LLM.” *Bird v. Microsoft Corp.*, No. 1:25-cv-05282 (S.D.N.Y. 2025) (ECF No. 7, at 2). Microsoft also accesses The Pile and Bibliotik, the “‘notorious pirated collection’ of ‘pirated books.’” *Id.* at 14.

133. Copyright holders filed an action against Apple in 2025 for “using unlicensed copyrighted works ... in building its artificial intelligence models.”²⁵ “Apple used Books3 to train its OpenELM language models [and] trained its Foundation Language Models using the same pirated dataset.” *Hendrix v. Apple Inc.*, No. 4:25-cv-07558-YGR (ECF No. 1, at 2).

134. Copyright holders also filed claims for copyright infringement based on accessing, scraping, and downloading files by circumventing YouTube’s access controls, including *Ted Entertainment, Inc. v. OpenAI, Inc.*, No. 5:26-cv-02935-BLF (N.D. Cal. 2026); and *Ted Entertainment, Inc. v. Amazon.com, Inc.*, 2:26-cv-01134 (W.D. Wash. 2026).

135. In addition, multiple BIPA lawsuits were filed in accord with the BIPA lawsuit filed against NVIDIA, including: *Basich v. Microsoft Corp.*, No. 2:26-cv-00422 (W.D. Wash. 2026); *Amer v. Eleven Labs Inc.*, No. 1:26-cv-05437 (N.D. Ill. 2026); *Marin v. Alphabet, Inc.*, No. 1:26-cv-05436 (N.D. Ill. 2026); *Marin v. Meta Platforms, Inc.*, No. 1:26-cv-05438 (N.D. Ill. 2026); *Lacour v. Apple Inc.*, No. 1:26-cv-05536 (N.D. Ill. 2026); *Dorcus v. Adobe Inc.*, No. 1:26-cv-05575 (N.D. Ill. 2026); and *Flowers v. Microsoft Corp.*,

²⁵ Rae Ann Varona, *Apple Using Pirated Books To Train AI Models, Authors Allege*, LAW360, Sept. 5, 2025.

No. 1:26-cv-05491 (N.D. Ill. 2026).

H. Materially False and Misleading Statements and Omissions in NVIDIA's SEC Filings and Code of Conduct

136. Defendants also made or approved false statements in various NVIDIA filings with the SEC, including misstatements and omissions in the Company's annual reports and annual proxy filings, and in the Company's Code of Conduct, all of which were approved by the Director Defendants.

i. False and Misleading Annual Reports

137. On February 24, 2023, NVIDIA filed its annual report on Form 10-K with the SEC for the fiscal year ended January 29, 2023 (the "2023 Annual Report"). The 2023 Annual Report was prepared and signed by all of the Director Defendants, as well as Defendants Robertson and Kress, which contained many false representations and material omissions regarding products based on the Company's AI platforms:

Fueled by the sustained demand for exceptional 3D graphics and the scale of the gaming market, NVIDIA has leveraged its GPU architecture to create platforms for scientific computing, artificial intelligence, or AI, data science, autonomous vehicles, or AV, robotics, metaverse and 3D internet applications.

The GPU was initially used to simulate human imagination, enabling the virtual worlds of video games and films. Today, it also simulates human intelligence, enabling a deeper understanding of the physical world. Its parallel processing capabilities, supported by thousands of computing cores, are essential to running deep learning algorithms. This form of AI, in which software writes itself by learning from large amounts of data, can serve as the brain of computers, robots and self-driving cars that can perceive and understand the world. GPU-powered deep learning is being adopted by thousands of enterprises to deliver services and products that would have been immensely difficult with traditional coding. *Some of the most recent applications of GPU-powered deep learning include recommendation systems, which are AI algorithms trained to understand the preferences, previous decisions, and characteristics of people and products using data gathered about their interactions, large language models, which can recognize, summarize, translate, predict and generate text and other content based on knowledge gained from massive datasets, and generative AI, which uses algorithms that create new content, including audio, code,*

images, text, simulations, and videos, based on the data they have been trained on.

NVIDIA has a platform strategy, bringing together hardware, systems, software, algorithms, libraries, and services to create unique value for the markets we serve. *While the computing requirements of these end markets are diverse, we address them with a unified underlying architecture leveraging our GPUs and software stacks. The programmable nature of our architecture allows us to support several multi-billion-dollar end markets with the same underlying technology by using a variety of software stacks developed either internally or by third-party developers and partners.* The large and growing number of developers across our platforms strengthens our ecosystem and increases the value of our platform to our customers.

Innovation is at our core. We have invested over \$37 billion in research and development since our inception, yielding inventions that are essential to modern computing. Our invention of the GPU in 1999 defined modern computer graphics and established NVIDIA as the leader in computer graphics. With our introduction of the CUDA programming model in 2006, we opened the parallel processing capabilities of our GPU for general purpose computing. This approach significantly accelerates the most demanding high-performance computing, or HPC, applications in fields such as aerospace, bio-science research, mechanical and fluid simulations, and energy exploration. Today, our GPUs and networking accelerate many of the fastest supercomputers across the world. In addition, the massively parallel computer architecture of our GPUs and associated software are well suited for deep learning and machine learning, powering the era of AI. *While traditional CPU-based approaches no longer deliver advances on the pace described by Moore's Law, NVIDIA accelerated computing delivers performance improvements on a pace ahead of Moore's Law, giving the industry a path forward.*

The NVIDIA computing platform is focused on accelerating the most compute-intensive workloads, such as AI, data analytics, graphics and scientific computing, across hyperscale, cloud, enterprise, public sector, and edge data centers. The platform consists of our energy efficient GPUs, data processing units, or DPUs, interconnects and systems, our CUDA programming model, and a growing body of software libraries, software development kits, or SDKs, application frameworks and services, which are either available as part of the platform or packaged and sold separately.

For both AI and HPC applications, the NVIDIA accelerated computing platform greatly increases computer and data center performance and power efficiency relative to conventional CPU-only approaches. In the field of AI, NVIDIA's platform accelerates both deep learning and machine learning workloads. Deep learning is a computer science approach where neural

networks are trained to recognize patterns from massive amounts of data in the form of images, sounds and text - in some instances better than humans — and in turn provide predictions in production use cases. Machine learning is a related approach that leverages algorithms as well as data to learn how to make determinations or predictions. HPC, which includes scientific computing, uses numerical computational approaches to solve large and complex problems.

We are engaged with thousands of organizations working on AI in a multitude of industries, from automating tasks such as consumer product and service recommendations, to chatbots for the automation of or assistance with live customer interactions, to enabling fraud detection in financial services, to optimizing oil exploration and drilling. These organizations include the world's leading consumer internet and cloud services companies, enterprises and startups seeking to implement AI in transformative ways across multiple industries. We partner with industry leaders to help transform their applications or their computing platforms. We also have partnerships in transportation, retail, healthcare, and manufacturing, among others, to accelerate the adoption of AI.

At the foundation of the NVIDIA accelerated computing platform are our GPUs, which excel at parallel workloads such as the training and inferencing of neural networks. They are available in industry standard servers from every major computer maker and CSP, as well as in our DGX AI supercomputer, a purpose-built system for deep learning and GPU accelerated applications. To facilitate customer adoption, we have also built other ready-to-use system reference designs around our GPUs, including HGX for hyperscale and supercomputing data centers, EGX for enterprise and edge computing, IGX for high-precision edge AI, and AGX for autonomous machines.

In addition to software that is delivered to customers as an integral part of our data center computing platform, we offer paid licenses to NVIDIA AI Enterprise, a comprehensive suite of enterprise-grade AI software; and NVIDIA vGPU software for graphics-rich virtual desktops and workstations.

Advancing the NVIDIA accelerated computing platform. NVIDIA's accelerated computing platform can solve complex problems in significantly less time and with lower power consumption than alternative computational approaches. Indeed, it can help solve problems that were previously deemed unsolvable. *We work to deliver continued performance leaps that outpace Moore's Law by leveraging innovation across the architecture, chip design, system, interconnect, and software layers. This full-stack innovation approach allows us to deliver order-of-magnitude*

performance advantages relative to legacy approaches in our target markets, which include Data Center, Gaming, Professional Visualization, and Automotive. While the computing requirements of these end markets are diverse, we address them with a unified underlying architecture leveraging our GPUs, CUDA and networking technologies as the fundamental building blocks. The programmable nature of our architecture allows us to make leveraged investments in research and development: we can support several multi-billion-dollar end markets with shared underlying technology by using a variety of software stacks developed either internally or by third-party developers and partners. We utilize this platform approach in each of our target markets.

Extending our technology and platform leadership in AI. We provide a complete, end-to-end accelerated computing platform for deep learning and machine learning, addressing both training and inferencing. This includes GPUs, interconnects, systems, our CUDA programming language, algorithms, libraries, and other software. GPUs are uniquely suited to AI, and we will continue to add AI-specific features to our GPU architecture to further extend our leadership position. Our AI technology leadership is reinforced by our large and expanding ecosystem in a virtuous cycle. Our GPU platforms are available from virtually every major server maker and CSP, as well as on our own AI supercomputer. There are 3.8 million developers worldwide using CUDA and our other software tools to help deploy our technology in our target markets. We evangelize AI through partnerships with hundreds of universities and over 13,000 startups through our Inception program. Additionally, our Deep Learning Institute provides instruction on the latest techniques on how to design, train, and deploy neural networks in applications using our accelerated computing platform.

Extending our technology and platform leadership in computer graphics. We believe that computer graphics is fundamental to the continued expansion and evolution of computing. *We apply our research and development resources to enhance the user experience for consumer entertainment and professional visualization applications, and create new virtual world and simulation capabilities. Our technologies are instrumental in driving gaming forward, as developers leverage our libraries and algorithms to deliver an optimized gaming experience on our GeForce platform. Our computer graphics platforms leverage not only our industry-leading GeForce and NVIDIA RTX GPUs, but also optimized software stacks. For example, GeForce Experience enhances each gamer's experience by optimizing their PC's settings, as well as enabling the recording and sharing of gameplay. Our Studio drivers enhance and accelerate a number of popular creative applications. Omniverse is real-time 3D design collaboration and virtual world simulation software that empowers artists, designers and creators to connect and collaborate in leading design applications.* We also enable interactive graphics applications - such as games, movie and photo editing and design software - to be accessed by almost any device, almost

anywhere, through our cloud platforms such as vGPU for enterprise and GeForce NOW for gaming. (Emphasis added.)

138. On February 21, 2024, NVIDIA filed its annual report on Form 10-K with the SEC for the fiscal year ended January 28, 2024 (the “2024 Annual Report”). The 2024 Annual Report was prepared and signed by all of the Director Defendants, as well as Defendants Robertson and Kress, and contained many false representations and material omissions regarding products based on the Company’s AI platforms:

Our full-stack includes the foundational CUDA programming model that runs on all NVIDIA GPUs, as well as hundreds of domain-specific software libraries, software development kits, or SDKs, and Application Programming Interfaces, or APIs. *This deep and broad software stack accelerates the performance and eases the deployment of NVIDIA accelerated computing for computationally intensive workloads such as artificial intelligence, or AI, model training and inference, data analytics, scientific computing, and 3D graphics, with vertical-specific optimizations to address industries ranging from healthcare and telecom to automotive and manufacturing.*

Our data-center-scale offerings are comprised of compute and networking solutions that can scale to tens of thousands of GPU-accelerated servers interconnected to function as a single giant computer; this type of data center architecture and scale is needed for the development and deployment of modern AI applications.

The GPU was initially used to simulate human imagination, enabling the virtual worlds of video games and films. Today, it also simulates human intelligence, enabling a deeper understanding of the physical world. Its parallel processing capabilities, supported by thousands of computing cores, are essential for deep learning algorithms. This form of AI, in which software writes itself by learning from large amounts of data, can serve as the brain of computers, robots and self-driving cars that can perceive and understand the world. GPU-powered AI solutions are being developed by thousands of enterprises to deliver services and products that would have been immensely difficult or even impossible with traditional coding. Examples include generative AI, which can create new content such as text, code, images, audio, video, and molecule structures, and recommendation systems, which can recommend highly relevant content such as products, services, media or ads using deep neural networks trained on vast datasets that capture the user preferences.

NVIDIA has a platform strategy, bringing together hardware, systems, software, algorithms, libraries, and services to create unique value for the

markets we serve. While the computing requirements of these end markets are diverse, we address them with a unified underlying architecture leveraging our GPUs and networking and software stacks. The programmable nature of our architecture allows us to support several multi-billion-dollar end markets with the same underlying technology by using a variety of software stacks developed either internally or by third-party developers and partners. The large and growing number of developers and installed base across our platforms strengthens our ecosystem and increases the value of our platform to our customers.

Innovation is at our core. We have invested over \$45.3 billion in research and development since our inception, yielding inventions that are essential to modern computing. Our invention of the GPU in 1999 sparked the growth of the PC gaming market and redefined computer graphics. With our introduction of the CUDA programming model in 2006, we opened the parallel processing capabilities of our GPU to a broad range of compute-intensive applications, paving the way for the emergence of modern AI. In 2012, the AlexNet neural network, trained on NVIDIA GPUs, won the ImageNet computer image recognition competition, marking the “Big Bang” moment of AI. We introduced our first Tensor Core GPU in 2017, built from the ground-up for the new era of AI, and our first autonomous driving system-on-chips, or SoC, in 2018. Our acquisition of Mellanox in 2020 expanded our innovation canvas to include networking and led to the introduction of a new processor class — the data processing unit, or DPU. Over the past 5 years, we have built full software stacks that run on top of our GPUs and CUDA to bring AI to the world’s largest industries, including NVIDIA DRIVE stack for autonomous driving, Clara for healthcare, and Omniverse for industrial digitalization; and introduced the NVIDIA AI Enterprise software — essentially an operating system for enterprise AI applications. In 2023, we introduced our first data center CPU, Grace, built for giant-scale AI and high-performance computing. With a strong engineering culture, we drive fast, yet harmonized, product and technology innovations in all dimensions of computing including silicon, systems, networking, software and algorithms. More than half of our engineers work on software.

The world’s leading cloud service providers, or CSPs, and consumer internet companies use our data center-scale accelerated computing platforms to enable, accelerate or enrich the services they deliver to billions of end users, including AI solutions and assistants, search, recommendations, social networking, online shopping, live video, and translation.

The NVIDIA Data Center platform is focused on accelerating the most compute-intensive workloads, such as AI, data analytics, graphics and scientific computing, delivering significantly better performance and power efficiency relative to conventional CPU-only approaches. It is

deployed in cloud, hyperscale, on-premises and edge data centers. The platform consists of compute and networking offerings typically delivered to customers as systems, subsystems, or modules, along with software and services.

Our compute offerings include supercomputing platforms and servers, bringing together our energy efficient GPUs, DPUs, interconnects, and fully optimized AI and high-performance computing, or HPC, software stacks. In addition, they include NVIDIA AI Enterprise software; our DGX Cloud service; and a growing body of acceleration libraries, APIs, SDKs, and domain-specific application frameworks.

Our networking offerings include end-to-end platforms for InfiniBand and Ethernet, consisting of network adapters, cables, DPUs, and switch systems, as well as a full software stack. This has enabled us to architect data center-scale computing platforms that can interconnect thousands of compute nodes with high-performance networking. While historically the server was the unit of computing, as AI and HPC workloads have become extremely large spanning thousands of compute nodes, the data center has become the new unit of computing, with networking as an integral part.

Our end customers include the world's leading public cloud and consumer internet companies, thousands of enterprises and startups, and public sector entities. We work with industry leaders to help build or transform their applications and data center infrastructure. Our direct customers include original equipment manufacturers, or OEMs, original device manufacturers, or ODMs, system integrators and distributors which we partner with to help bring our products to market. We also have partnerships in automotive, healthcare, financial services, manufacturing, and retail among others, to accelerate the adoption of AI.

At the foundation of the NVIDIA accelerated computing platform are our GPUs, which excel at parallel workloads such as the training and inferencing of neural networks. They are available in the NVIDIA accelerated computing platform and in industry standard servers from every major cloud provider and server maker. Beyond GPUs, our data center platform expanded to include DPUs in fiscal year 2022 and CPUs in fiscal year 2024. We can optimize across the entire computing, networking and storage stack to deliver data center-scale computing solutions.

While our approach starts with powerful chips, what makes it a full-stack computing platform is our large body of software, including the CUDA parallel programming model, the CUDA-X collection of acceleration libraries, APIs, SDKs, and domain-specific application frameworks.

In addition to software delivered to customers as an integral part of our data center computing platform, we offer paid licenses to NVIDIA AI Enterprise, a comprehensive suite of enterprise-grade AI software and

NVIDIA vGPU software for graphics-rich virtual desktops and workstations.

In fiscal year 2024, we launched the NVIDIA DGX Cloud, an AI-training-as-a-service platform which includes cloud-based infrastructure and software for AI, customizable pretrained AI models, and access to NVIDIA experts. We have partnered with leading cloud service providers to host this service in their data centers.

The programmable nature of our architecture allows us to make leveraged investments in research and development: *we can support several multi-billion-dollar end markets with shared underlying technology by using a variety of software stacks developed either internally or by third-party developers and partners. We utilize this platform approach in each of our target markets.*

Extending our technology and platform leadership in AI. We provide a complete, end-to-end accelerated computing platform for AI, addressing both training and inferencing. This includes full-stack data center-scale compute and networking solutions across processing units, interconnects, systems, and software. Our compute solutions include all three major processing units in AI servers — GPUs, CPUs, and DPUs. GPUs are uniquely suited to AI, and we will continue to add AI-specific features to our GPU architecture to further extend our leadership position. In addition, we offer DGX Cloud, an AI-training-as-a-service platform, and NeMo – a complete solution for building enterprise-ready Large Language Models, or LLMs, using open source and proprietary LLMs created by NVIDIA and third parties. Our AI technology leadership is reinforced by our large and expanding ecosystem in a virtuous cycle. Our computing platforms are available from virtually every major server maker and CSP, as well as on our own AI supercomputers. There are over 4.7 million developers worldwide using CUDA and our other software tools to help deploy our technology in our target markets. *We evangelize AI through partnerships with hundreds of universities and thousands of startups through our Inception program. Additionally, our Deep Learning Institute provides instruction on the latest techniques on how to design, train, and deploy neural networks in applications using our accelerated computing platform.*

Extending our technology and platform leadership in computer graphics. *We believe that computer graphics infused with AI is fundamental to the continued expansion and evolution of computing. We apply our research and development resources to enhance the user experience for consumer entertainment and professional visualization applications and create new virtual world and simulation capabilities. Our technologies are instrumental in driving the gaming, design, and creative industries forward, as developers leverage our libraries and algorithms to*

deliver an optimized experience on our GeForce and NVIDIA RTX platforms. Our computer graphics platforms leverage AI end-to-end, from the developer tools and cloud services to the Tensor Cores included in all RTX-class GPUs. For example, NVIDIA Avatar Cloud Engine, or ACE, is a suite of technologies that help developers bring digital avatars to life with generative AI, running in the cloud or locally on the PC.

Concerns relating to the responsible use of new and evolving technologies, such as AI, in our products and services may result in reputational or financial harm and liability and may cause us to incur costs to resolve such issues. *We are increasingly building AI capabilities and protections into many of our products and services, and we also offer stand-alone AI applications. AI poses emerging legal, social, and ethical issues and presents risks and challenges that could affect its adoption, and therefore our business. If we enable or offer solutions that draw controversy due to their perceived or actual impact on society, such as AI solutions that have unintended consequences, infringe copyright or rights of publicity, or are controversial because of their impact on human rights, privacy, employment or other social, economic or political issues, or if we are unable to develop effective internal policies and frameworks relating to the responsible development and use of AI models and systems offered through our sales channels, we may experience brand or reputational harm, competitive harm or legal liability. Complying with multiple regulations from different jurisdictions related to AI could increase our cost of doing business, may change the way that we operate in certain jurisdictions, or may impede our ability to offer certain products and services in certain jurisdictions if we are unable to comply with regulations. Compliance with existing and proposed government regulation of AI, including in jurisdictions such as the European Union as well as under any U.S. regulation adopted in response to the Biden administration's Executive Order on AI, may also increase the cost of related research and development, and create additional reporting and/or transparency requirements.* For example, regulation adopted in response to the Executive Order on AI could require us to notify the USG of certain safety test results and other information. Furthermore, changes in AI-related regulation could disproportionately impact and disadvantage us and require us to change our business practices, which may negatively impact our financial results. *Our failure to adequately address concerns and regulations relating to the responsible use of AI by us or others could undermine public confidence in AI and slow adoption of AI in our products and services or cause reputational or financial harm.* (Emphasis added.)

139. On February 26, 2025, NVIDIA filed its annual report on Form 10-K with the SEC for the fiscal year ended January 26, 2025 (the “2025 Annual Report”). The 2025

Annual Report was prepared and signed by all the Director Defendants, as well as Defendants Robertson and Kress, and contained many false representations, and material omissions regarding products based on the Company's AI platforms:

Our full-stack includes the foundational CUDA programming model that runs on all NVIDIA GPUs, as well as hundreds of domain-specific software libraries, software development kits, or SDKs, and Application Programming Interfaces, or APIs. This deep and broad software stack accelerates the performance and eases the deployment of NVIDIA accelerated computing for computationally intensive workloads such as artificial intelligence, or AI, model training and inference, data analytics, scientific computing, and 3D graphics, with vertical-specific optimizations to address industries ranging from healthcare and telecom to automotive and manufacturing.

Our data-center-scale offerings are comprised of compute and networking solutions that can scale to tens of thousands of GPU-accelerated servers interconnected to function as a single giant computer; this type of data center architecture and scale is needed for the development and deployment of modern AI applications.

The GPU was initially used to simulate human imagination, enabling the virtual worlds of video games and films. Today, it also simulates human intelligence, enabling a deeper understanding of the physical world. Its parallel processing capabilities, supported by thousands of computing cores, are essential for deep learning algorithms. This form of AI, in which software writes itself by learning from large amounts of data, can serve as the brain of computers, robots, and self-driving cars that can perceive and understand the world. GPU-powered AI solutions are being developed by thousands of enterprises to deliver services and products that would have been immensely difficult or even impossible with traditional coding. Examples include generative AI, which can create new content such as text, code, images, audio, video, molecule structures, and recommendation systems, which can recommend highly relevant content such as products, services, media, or ads using deep neural networks trained on vast datasets that capture the user's preferences.

NVIDIA has a platform strategy, bringing together hardware, systems, software, algorithms, libraries, and services to create unique value for the markets we serve. While the computing requirements of these end markets are diverse, we address them with a unified underlying architecture leveraging our GPUs and networking and software stacks. The programmable nature of our architecture allows us to support several multi-billion-dollar end markets with the same underlying technology by using a variety of software stacks developed either internally or by third-party developers and partners. The large and growing number of developers and

installed base across our platforms strengthens our ecosystem and increases the value of our platform to our customers.

Innovation is at our core. We have invested over \$58.2 billion in research and development since our inception, yielding inventions that are essential to modern computing. Our invention of the GPU in 1999 sparked the growth of the PC gaming market and redefined computer graphics. With our introduction of the CUDA programming model in 2006, we opened the parallel processing capabilities of our GPU to a broad range of compute-intensive applications, paving the way for the emergence of modern AI. In 2012, the AlexNet neural network, trained on NVIDIA GPUs, won the ImageNet computer image recognition competition, marking the “Big Bang” moment of AI. We introduced our first Tensor Core GPU in 2017, built from the ground-up for the new era of AI, and our first autonomous driving system-on-chips, or SoC, in 2018. Our acquisition of Mellanox in 2020 expanded our innovation canvas to include networking, enabled our platforms to be data center scale, and led to the introduction of a new processor class – the data processing unit, or DPU. Over the past 5 years, we have built full software stacks that run on top of our GPUs and CUDA to bring AI to the world’s largest industries, including NVIDIA DRIVE stack for autonomous driving, Clara for healthcare, and Omniverse for industrial digitalization; and introduced the NVIDIA AI Enterprise software – essentially an operating system for enterprise AI applications. ***In 2023, we introduced our first data center CPU, Grace, built for giant-scale AI and high performance computing, or HPC. With a strong engineering culture, we drive fast, yet harmonized, product and technology innovations in all dimensions of computing including silicon, systems, networking, software and algorithms. More than half of our engineers work on software.***

The world’s leading cloud service providers, or CSPs, and consumer internet companies use our data center-scale accelerated computing platforms to enable, accelerate, develop, or enrich the services and offerings they deliver to billions of end users, including AI solutions and assistants, AI foundation models, search, recommendations, social networking, online shopping, live video, and translation.

Enterprises and startups across a broad range of industries use our accelerated computing platforms to build new generative and agentic AI-enabled products and services, and/or to dramatically accelerate and reduce the costs of their workloads and workflows. The enterprise software industry uses them for new AI assistants, chatbots, and agents; the transportation industry for autonomous driving; the healthcare industry for accelerated and computer-aided drug discovery; and the financial services industry for customer support and fraud detection.

The NVIDIA Data Center platform is focused on accelerating the most compute-intensive workloads, such as AI, data analytics, graphics, and scientific computing, delivering significantly better performance and power efficiency relative to conventional CPU-only approaches. It is deployed in cloud, hyperscale, on-premises and edge data centers. ***The platform consists of compute and networking offerings typically delivered to customers as systems, subsystems, or modules, along with software and services.***

Our compute offerings include supercomputing platforms and servers, bringing together our energy efficient GPUs, CPUs, interconnects, and fully optimized AI and HPC software stacks. In addition, they include NVIDIA AI Enterprise software; our DGX Cloud service; and a growing body of acceleration libraries, APIs, SDKs, and domain-specific application frameworks.

Our networking offerings include end-to-end platforms for InfiniBand and Ethernet, consisting of network adapters, cables, DPUs, switch chips and systems, as well as a full software stack. This has enabled us to architect data center-scale computing platforms that can interconnect thousands of compute nodes with high-performance networking. While historically the server was the unit of computing, as AI and HPC workloads have become extremely large spanning thousands of compute nodes, the data center has become the new unit of computing, with networking as an integral part.

Our customers include the world's leading public cloud and consumer internet companies, thousands of enterprises and startups, and public sector entities. We work with industry leaders to help build or transform their applications and data center infrastructure. Our direct customers include original equipment manufacturers, or OEMs, original device manufacturers, or ODMs, system integrators and distributors which we partner with to help bring our products to market. We also have partnerships in automotive, healthcare, financial services, manufacturing, retail, and technology among others, to accelerate the adoption of AI.

At the foundation of the NVIDIA accelerated computing platform are our GPUs, which excel at parallel workloads such as the training and inferencing of neural networks. They are available in the NVIDIA accelerated computing platform and in industry standard servers from every major cloud provider and server maker. Beyond GPUs, our data center platform expanded to include DPUs in fiscal year 2022 and CPUs in fiscal year 2024. We can optimize across the entire computing, networking and storage stack to deliver data center-scale computing solutions.

While our approach starts with powerful chips, what makes it a full-stack computing platform is our large body of software, including the CUDA parallel programming model, the CUDA-X collection of acceleration libraries, APIs, SDKs, and domain-specific application frameworks.

In addition to software delivered to customers as an integral part of our data center computing platform, we offer paid licenses to NVIDIA AI Enterprise, a comprehensive suite of enterprise-grade AI software and NVIDIA vGPU software for graphics-rich virtual desktops and workstations. We also offer the NVIDIA DGX Cloud, a fully managed AI-training-as-a-service platform which includes cloud-based infrastructure and software for AI, customizable pretrained AI models, and access to NVIDIA experts.

In fiscal year 2025, we launched the NVIDIA Blackwell architecture, a full set of data center scale infrastructure that includes GPUs, CPUs, DPUs, interconnects, switch chips and systems, and networking adapters. Blackwell excels at processing cutting edge generative AI and accelerated computing workloads with market leading performance and efficiency. Offered in a number of configurations, it can address the needs of customers across industries and a diverse set of AI and accelerated computing use cases.

Our accelerated computing platform can solve complex problems in significantly less time and with lower power consumption than alternative computational approaches. Indeed, it can help solve problems that were previously deemed unsolvable. We work to deliver continued performance leaps that outpace Moore's Law by leveraging innovation across the architecture, chip design, system, interconnect, algorithm, and software layers. This full-stack innovation approach allows us to deliver order-of-magnitude performance advantages relative to legacy approaches in our target markets, which include Data Center, Gaming, Professional Visualization, and Automotive. While the computing requirements of these end markets are diverse, we address them with a unified underlying architecture leveraging our GPUs, CPUs, CUDA and networking technologies as the fundamental building blocks. The programmable nature of our architecture allows us to make leveraged investments in research and development: we can support several multi-billion-dollar end markets with shared underlying technology by using a variety of software stacks developed either internally or by third-party developers and partners. We utilize this platform approach in each of our target markets.

We provide a complete, end-to-end accelerated computing platform for AI, addressing both training and inferencing. This includes full-stack data center-scale compute and networking solutions across processing units, interconnects, systems, and software. *Our compute solutions include all three major processing units in AI servers — GPUs, CPUs, and DPUs. GPUs are uniquely suited to AI, and we will continue to add AI-specific features to our GPU architecture to further extend our leadership*

position.

In addition, we offer DGX Cloud, a fully managed AI-training-as-a-service platform, along with NVIDIA AI Enterprise—a comprehensive software suite designed to simplify the development and deployment of production-grade, end-to-end generative AI applications. NVIDIA AI Enterprise includes: NVIDIA NIM, which delivers a 2.5x increase in token throughput using industry-leading open and proprietary models; NVIDIA NeMo, a complete solution for curating, fine-tuning, evaluating, and safeguarding domain-adapted models; and AI Blueprints, pre-built, runnable templates that help enterprises build, optimize, and deploy AI agents while preserving privacy. These tools enable organizations to securely develop and run AI applications on NVIDIA-accelerated infrastructure anywhere.

Our AI technology leadership is reinforced by our large and expanding ecosystem in a virtuous cycle. Our computing platforms are available from virtually every major server maker and CSP, as well as on our own AI supercomputers. There are over 5.9 million developers worldwide using CUDA and our other software tools to help deploy our technology in our target markets. We evangelize AI through partnerships with hundreds of universities and thousands of startups through our Inception program. Additionally, our Deep Learning Institute provides instruction on the latest techniques on how to design, train, and deploy neural networks in applications using our accelerated computing platform.

We believe that computer graphics infused with AI is fundamental to the continued expansion and evolution of computing. We apply our research and development resources to enhance the user experience for consumer entertainment and professional visualization applications and create new virtual world and simulation capabilities. Our technologies are instrumental in driving the gaming, design, and creative industries forward, as developers leverage our libraries and algorithms to deliver an optimized experience on our GeForce and NVIDIA RTX platforms. Our computer graphics platforms leverage AI end-to-end, from the developer tools and cloud services to the Tensor Cores included in all RTX-class GPUs. For example, NVIDIA Avatar Cloud Engine, or ACE, is a suite of technologies that help developers bring digital avatars to life with generative AI, running in the cloud or locally on the PC. GeForce Experience enhances each gamer's experience by optimizing their PC's settings, as well as enabling the recording and sharing of gameplay. Our Studio drivers enhance and accelerate a number of popular creative applications. Omniverse is real-time 3D design collaboration and virtual world simulation software that empowers artists, designers, and creators to connect and collaborate in leading design applications. We also enable interactive graphics applications - such as games, movie and photo editing and design software - to be accessed by almost any device, almost

anywhere, through our cloud platforms such as vGPU for enterprise and GeForce NOW for gaming.

ii. False and Misleading Proxy Statements

140. On May 8, 2023, NVIDIA filed its annual Proxy Statement on Schedule 14A with the SEC (the “2023 Proxy Statement”). The 2023 Proxy Statement was reviewed and authorized by each of the Director Defendants and was filed in connection with the 2023 stockholder meeting to elect all of the Director Defendants, among other things.

141. The Director Defendants caused false statements and material omissions to be included in the 2023 Proxy Statement, which stated that NVIDIA was focused on “information security and privacy protections” to establish AI services. The Proxy included specific representations regarding the data used to train NVIDIA’s AI services, evidencing the importance of NVIDIA’s AI development and the Director Defendants’ knowledge thereof:

We design our products to protect the privacy, networks, computers, programs, information and data of our customers, partners, and employees. The Board is committed to strong and meaningful information security and privacy protections. Our Chief Security Officer and members of our security team present at least annually to our Board and provide updates throughout the year as needed. These leaders also update the AC quarterly.

Our information security, including cybersecurity, practices comprise the physical, procedural, and technical safeguards we take and are designed to protect customer and employee information from unauthorized access or attack, and measures designed to secure NVIDIA networks, systems, devices, products, and services in order to secure the privacy of our customers’ and employees’ data. We established a cross-functional leadership team, consisting of executive-level leaders, that meets monthly to review cybersecurity matters and evaluate emerging threats. To ensure a robust breadth of knowledge, the team consults as needed with external parties, such as computer security firms and risk management and governance experts. With oversight and guidance provided by the cross-functional leadership team, our information security teams continually refine our practices to address emerging security risks and changes in regulations.

We have a privacy policy that describes how we collect, use, store, process, share and protect customer data, as well as how customers can access and manage their personal data. We seek to uphold the legal protections safeguarding the privacy of our customers' data. Our employees are required to complete information security awareness training and to comply with our information security and privacy policies.

We seek to advance trustworthy AI that is founded in our core values, reflects our Code of Conduct and is rooted in the principles of upholding human rights. We recognize that technology can have a profound impact on people and the world and have therefore set priorities that aim to foster positive change and enable trust and transparency in AI development.

Our products are programmable and general purpose in nature. When we provide tools to help developers create applications for specific industries, we focus on creating products and services that enable developers to create and accelerate socially beneficial applications.

(Emphasis added.)

142. On May 14, 2024, NVIDIA filed its annual Proxy Statement on report on Schedule 14A with the SEC (the "2024 Proxy Statement"). The 2024 Proxy Statement was reviewed and authorized by each of the Director Defendants and was filed in connection with the 2024 stockholder meeting to elect all of the Director Defendants, among other things.

143. The Director Defendants in the 2024 Proxy Statement largely repeated their prior representations that the Company's AI training dataset contained exclusively public domain materials, made affirmative statements and omitted material facts that NVIDIA was using extensive collections of pirated works to train the Company's AI software:

Our trustworthy artificial intelligence, or AI, principles, which we share with customers and partners, reflect our core values and our Code of Conduct. ***We endeavor to deliver AI models that comply with privacy and data protection laws, perform safely and as intended, provide transparency about a model's design and limitations, minimize unwanted bias, and give equal opportunity to benefit from AI.***

Our products are programmable and general purpose in nature. When we provide tools to help developers create applications for specific industries,

we focus on creating products and services that enable developers to create and accelerate socially beneficial applications.

Our NCGC oversees our public policy engagement and accountability. Our Government Relations team engages in public policy advocacy to affect government action on issues of importance to our business, customers, stockholders, and employees, and to provide thought leadership to global governments on issues that directly affect our business. It is also a platform for educating policymakers through demonstrations of NVIDIA's technology, amplifying our work in targeted areas, and collaborating with various organizations on issues of shared interest. We focus our public policy activities in AI, specifically to promote investment in core AI research, support workforce development around AI, and provide educational resources to technology policy advisors. NVIDIA may incur expenditures to support or educate viewpoints on public policy issues, including expenditures for intermediaries that advocate on our behalf if it is in our best interest.

144. On May 13, 2025, NVIDIA filed its annual Proxy Statement on report on Schedule 14A with the SEC (the "2025 Proxy Statement"). The 2025 Proxy Statement was authorized by each of the Director Defendants and was filed in connection with the 2025 stockholder meeting to elect all the Director Defendants, among other things.

145. The Director Defendants in the 2025 Proxy Statement repeated their many false representations and material omissions that the dataset contained exclusively public domain materials:

Our artificial intelligence, or AI, principles, which we share with customers and partners, reflect our core values and our Code of Conduct. We endeavor to deliver trustworthy AI models that comply with privacy and data protection laws, perform safely and as intended, provide transparency about a model's design and limitations, minimize unwanted bias, and give equal opportunity to benefit from AI.

Our products are programmable and non-specific in nature. When we provide tools to help developers create applications for specific industries, we focus on creating products and services that enable developers to create and accelerate socially beneficial applications.

146. These statements were materially false and misleading because they led the market and the Company's shareholders to believe that NVIDIA had complied with all laws, including copyright laws, in the development of its AI. In a similar vein, the statements in the Proxy contained material omissions because the Proxy failed to disclose the material fact that NVIDIA had trained its AI using vast amounts of data for which it had not obtained the necessary licenses or permissions, and thus that the Company was at material risk of lawsuits from copyright holders. These statements and omissions were made at a time when all Defendants knew or should have known that other AI developers had been sued by copyright holders for using datasets that included unlicensed materials, and that the datasets used by NVIDIA were not subjected to "principles [that] reflect our core values and our Code of Conduct," and that Defendants did not "endeavor to deliver trustworthy AI models that comply with privacy and data protection laws [or] perform safely and as intended." At best, the Director Defendants acted recklessly in making representations regarding NVIDIA's practices given the copyright claims associated with The Pile, Books3, Anna's Archive, RedPajama, LibGen, SciHub, Z-Library, and Panda-70M, datasets comprised of YouTube videos, and human speech recordings.

147. On May 12, 2026, NVIDIA filed its annual Proxy Statement on report on Schedule 14A with the SEC (the "2026 Proxy Statement"). The 2026 Proxy Statement, again reviewed and authorized by each of the Director Defendants, conspicuously omitted the representations made in the 2023 Proxy Statement, the 2024 Proxy Statement, and the 2025 Proxy Statement that the data used by NVIDIA were subjected to "principles [that] reflect our core values and our Code of Conduct," and that Defendants "endeavor to deliver trustworthy AI models that comply with privacy and data protection laws [or] perform safely and as intended."

148. The Director Defendants’ effective retraction of their prior representations regarding datasets used to develop their AI models is a tacit admission that the datasets were decidedly not free from copyrighted materials at the time of the 2023, 2024, and 2025 Proxy Statements.

iii. False and Misleading Code of Conduct

149. Defendants referred to NVIDIA’s Code of Conduct, available on the NVIDIA website, in each of their Annual Reports and Proxy Statements referred to above. The Court has held in multiple actions that material misrepresentations in a corporation’s code of business conduct on a corporation’s website can be actionable under § 10(b).²⁶

150. NVIDIA’s Code of Conduct contains material misrepresentations that mislead NVIDIA stockholders to believe Defendants did not cause NVIDIA to engage in copyright infringement:

(a) “NVIDIA pioneers accelerated AI computing innovating across the entire stack, the whole data center, and from cloud to robotic systems.” “We prioritize developing trustworthy AI to ensure this powerful technology provides benefits.” “We value diversity and inclusion, which enables us to achieve this goal. Additionally, we recognize our responsibility to maintain high ethical standards and identify concerns”

(u) “*We don’t engage in activities that might limit competition or violate antitrust law.*” “We gather competitive information with care, seeking only data that is publicly available or licensed to

²⁶ See *Flynn v. Exelon Corp.*, 2021 WL 1561712 (N.D. Ill. Apr. 21, 2021); *Holwill v. AbbVie Inc.*, 2020 WL 5235005 (N.D. Ill. Sept. 1, 2020).

us.” *“We’re committed to full, fair, accurate, timely, and clear disclosures in reports and documents that we file with or submit to government agencies, as well as in other public communications or in material that we develop for internal use.”*

- (v) “We support the use of our products for socially beneficial purposes and work to address the risks of potential harm.” *“We ensure our customers can trust our products by sharing important information about how they are built and work.”*
“When it comes to artificial intelligence, we take a proactive approach to identity and address issues of trustworthiness concerns through the development process—from data collection to decommissioning. When necessary, our AI Ethics Committee considers the implications of operating product.”
- (w) “We don’t create misleading impressions in any advertising, marketing, or sales materials or presentations and don’t make false or illegal claims about competitors or their offerings.”
- (x) “For our artificial intelligence products, we use model cards to explain how a model was designed, trained, and tested, how it performs, and its intended purposes.”

151. These statements were materially false and misleading and contained material omissions. The statements and omissions were made during the time frame when

all Defendants knew or should have known the datasets used by NVIDIA were not subjected to the principles set forth in the Code of Conduct,” and that, “[w]hen it comes to artificial intelligence,” Defendants did not “take a proactive approach to identify and address issues of trustworthiness concerns through the development process—from data collection to decommissioning.” Since Defendants chose to affirmatively represent that they were not violating competition and antitrust laws in the development of the Company’s AI, they had a duty to speak the whole truth and disclose that they did not have the necessary licenses or permission. This they recklessly disregarded.

I. Share Repurchases

152. During the same time period that NVIDIA was developing its AI platforms based on the use of unauthorized use of copyrighted works in developing NVIDIA’s LLMs, VLMs, and CVMs, the Board, which was directly responsible for the Company’s stock repurchase program, caused NVIDIA to repurchase its own common stock at prices that were artificially inflated by the materially false and misleading statements and omissions in NVIDIA’s SEC filings. The Board authorized multiple share repurchase programs in recent fiscal years, while aware of the true extent of the Company’s copyright and BIPA infringement on its AI platforms and datasets, and the false and misleading Proxy Statements and Annual Reports. In May 2022, the Board authorized a share repurchase program up to \$15 billion through December 2023. In August 2023, the Board authorized a share repurchase program up to \$25 billion, without expiration. On August 26, 2024, the Board approved an additional share repurchase program of up to \$50 billion. In addition, on August 26, 2025, the Board approved an additional share repurchase program of up to \$60 billion.

153. As of January 28, 2024, just prior to the filing of the *Nazemian* action,

NVIDIA had repurchased approximately 12.3 million shares under this program: (i) 7 million shares purchased at an average price of \$148.11 per share between October 31, 2022 to November 27, 2022; (ii) 900,000 shares purchased at an average price of \$464.39 per share between October 30, 2023 to November 26, 2023; (iii) 1.1 million shares purchased at an average price of \$477.26 per share between November 27, 2023 to December 24, 2023; and (iv) 3.3 million shares purchased at an average price of \$540.85 per share between December 25, 2023 to January 28, 2024.

154. Subsequent to NVIDIA's share repurchase as of January 28, 2024, NVIDIA's stock began to astronomically inflate, from \$624.70 on January 29, 2024, to \$1208.90 on June 7, 2024. On June 10, 2024, NVIDIA completed a 10-for-1 stock split, which immediately deflated NVIDIA stock price down to \$121.79. As of January 26, 2025, just after the filing of the *Ted Entertainment* action, and just prior to the filing of the *Rogers* action, NVIDIA had repurchased approximately 74.8 million shares under this program: (i) 25.4 million shares purchased at an average price of \$142.67 per share between October 28, 2024 to November 24, 2024; (ii) 10.6 million shares purchased at an average price of \$136.86 per share between November 25, 2024 to December 22, 2024; (iii) 19.3 million shares purchased at an average price of \$39.30 per share between December 23, 2024 to January 26, 2025; (iv) 5.6 million shares purchased at an average price of \$198.89 per share between October 27, 2025 to November 23, 2025; (v) 6.6 million shares purchased at the average price of \$179.41 per share between November 24, 2025 to December 21, 2025; and (vi) 7.3 million shares purchased at the average price of \$186.52 per share between December 22, 2025 to January 25, 2026.

155. All of these repurchases were made at times when the Defendants knew, but had failed to disclose, the Company's copyright and BIPA infringement on its AI platforms

and datasets. As a result, the Company paid artificially inflated prices to buy its own shares.

J. Loss Causation/Economic Loss

156. During the Relevant Time Period, as detailed herein, the false statements and material omissions approved by the Individual Defendants deceived the market and artificially inflated the prices of NVIDIA common stock and operated as a fraud or deceit on purchasers or acquirers of NVIDIA common stock, including NVIDIA which was repurchasing its own stock at Defendants' direction. When Defendants' prior misrepresentations and fraudulent conduct were disclosed and became apparent to the market, the price of NVIDIA common stock fell precipitously. As a result of its purchases or acquisitions of NVIDIA common stock during the Relevant Time Period, NVIDIA suffered economic loss, *i.e.*, damages, under the federal securities laws.

157. The Individual Defendants' false and misleading statements and omissions had the intended effect and caused NVIDIA common stock to trade at artificially inflated levels throughout the time of NVIDIA's stock purchases.

158. On May 5, 2026, the date that the Court denied in part NVIDIA's motion to dismiss in the *Nazemian* action, NVIDIA stock dropped to \$196.50 per share from NVIDIA stock's highest trading date on April 27, 2026, at \$216.61 per share.

159. NVIDIA's stock fell on each trading date between May 20, 2026 and May 27, 2026 and, beginning on June 5, 2026, NVIDIA stock declined even further. As of June 26, 2026, NVIDIA's stock closed at \$192.53. This occurred despite the fact that, on May 20, 2026, NVIDIA announced record Q1 FY 2027 earnings of \$81.6 billion, up 20% from the previous quarter and up 85% from the previous year. This same date, NVIDIA also announced an \$80 billion additional share repurchase authorization.

160. The declines in the price of NVIDIA common stock after the disclosures came to light were a direct result of the nature and extent of the Individual Defendants' fraud finally being revealed to investors and the market and a materialization of the undisclosed risks alleged herein. The timing and magnitude of the price declines in NVIDIA common stock negates any inference that the loss suffered by NVIDIA was caused by changed market conditions, macroeconomic or industry factors or Company-specific facts unrelated to the Individual Defendants' fraudulent conduct. The economic loss, *i.e.*, damages, suffered by NVIDIA was a direct result of the Individual Defendants' fraudulent scheme to artificially inflate the price of NVIDIA common stock and the subsequent significant declines in the value of NVIDIA common stock when the Individual Defendants' prior misrepresentations and other fraudulent conduct were revealed.

K. The Individual Defendants Made Their Misrepresentations Knowingly or With Reckless Disregard for the Truth

161. When the Officer and Director Defendants made the public statements detailed above, they knew, or with extreme recklessness disregarded, that those statements were false and misleading, including by omitting material facts. The Individual Defendants knew their statements would be issued and disseminated to the investing public, knew analysts and investors were likely to rely upon those misrepresentations and omissions, and knowingly and recklessly participated in the issuance and dissemination of those statements. Indeed, the ongoing fraud as detailed herein could not have been perpetrated without the knowledge or recklessness of personnel at the highest level of the Company, including the Officer and Director Defendants.

162. The Individual Defendants acted with scienter in that Defendants knew, or recklessly disregarded, that the public statements they issued or disseminated during the relevant time period were false or misleading. Alternatively, the Individual Defendants

consciously or recklessly ignored red flags which demonstrated the falsity of Defendants' statements.

163. The facts, when viewed holistically and together with the other allegations in the Complaint, establish a strong inference of scienter that each of the 10(b) Defendants knew or was severely reckless in not knowing that each of the alleged misrepresentations and omissions was false and misleading at the time that it was made.

164. The Individual Defendants' scienter is demonstrated by the fact that they were actively involved in the development of NVIDIA's AI strategy, which was a key Board issue that had been prioritized as a "generational" opportunity.

165. The Individual Defendants were also aware of news articles at the time alleging copyright infringement and BIPA by NVIDIA. On March 11, 2024, a Law 360 article discussed how "Writers Target Nvidia, Databricks in AI Copyright Litigation." The article stated that NVIDIA was "accused of copyright infringements" by plaintiffs that "said they own registered copyrights in certain books that were included in the training dataset that Nvidia has admitted copying to train its NeMo Megatron models."²⁷ On January 21, 2026, the Individual Defendants were made aware that the Court agreed to permit these plaintiffs to amend their complaint—based on discovery in the *Nazemian* action—in an article entitled "NVIDIA must face direct infringement claim over 'shadow library' downloads."²⁸ The article informed the Individual Defendants that plaintiffs had "newly discovered AI models allegedly trained on the illegally downloaded works" and

²⁷ Adam Lidgett, *Writers Target Nvidia, Databricks in AI Copyright Litigation*, LAW360, Mar. 11, 2024.

²⁸ Thomas Long, *NVIDIA must face direct infringement claim over 'shadow library' downloads*, VITALLAW, Jan. 21, 2026.

produced “new claims of secondary infringement based on NVIDIA’s alleged acts encouraging its customers to download and use its training dataset.”

166. On December 1, 2025, the Individual Defendants were informed that “Nvidia [has been] sued for allegedly scraping copyright-protected video from YouTube, in potential class action.”²⁹ The Piracy Monitor described the *Ted Entertainment* lawsuit, where “[t]he plaintiffs contend that nvidia[] violates the provisions of the Digital Millenium Copyright Act (DMCA) regarding anti-circumvention ... by bypassing technological protection measures that control access to copying of YouTube videos.” The Piracy Monitor news article explained:

What is the content in question?

All three plaintiffs contend that they have invested substantial time and money into producing their works and bringing awareness around it

- Ted Entertainment, Inc. is an independent content creator which owns and operates YouTube channels “h3h3 Productions” and “H3 Podcast Highlights,” with more than 5,800 original videos on YouTube, with a combined total of over 4,000,000,000 YouTube views and more than 2,600,000 subscribers.
- Plaintiff Matt Fisher is an individual creator and California resident who posts golf videos on YouTube, many of which are instructional. His Mr.ShortGame YouTube channel has over 500,000 subscribers and hundreds of millions of views.
- Golfholics is a California corporation with a golf channel that has 130,000 subscribers and has had millions of views.

Intentional infringement?

nVidia was said to use a video-downloading program combined with virtual machines that rotated IP addresses to avoid detection and blocking; which enabled the “mass extraction of, at times, eighty years of video content per day.”

²⁹ *nVidia sued for allegedly scraping copyright-protected video from YouTube, in potential class action*, PIRACY MONITOR, Dec. 1, 2025.

Internal conversations detailed in the Complaint demonstrate that nVidia personnel were aware of the protections and actively worked toward circumventing them.

167. On May 13, 2026, the Individual Defendants were informed of the *Rogers* action, “Illinois suits accuse Microsoft, Nvidia of stealing voiceprints for AI products.” In the Courthouse News Service article, Destiny Devooght explained that “Seven Illinois professionals claim Microsoft and Nvidia stole their unique voiceprints and used them to train commercial artificial intelligence voice synthesis models”:

The plaintiffs claim Nvidia ... obtained recordings of their voices, extracted the unique biometric signatures from the stolen voiceprints and used the data to build suites of commercial voice artificial intelligence products.

A voiceprint is the digital fingerprint of a human voice — a mathematical representation of its pitch, timbre, resonance, accent and distinctive speech patterns.

The companies used the voiceprints of highly accomplished professionals without consent or notification, according to the plaintiffs, knowingly violating Illinois’ Biometric Information Privacy Act for the construction of products like ... Nvidia’s Canary and Parakeet speech models.

The act requires written notice of the specific purpose and duration of voiceprint collection, along with a written release from the speaker.

Each suite of AI voice products utilized hundreds of thousands of hours of recordings from just as many speakers, including from each plaintiff, they say. Collection and training at that scale includes major compliance responsibilities, including identifying source speakers, providing notice and obtaining and storing countless releases.

168. The Individual Defendants were aware of these public concessions of infringement, and thus had a duty to investigate further, especially in light of the enormous financial risks to NVIDIA of infringement lawsuits.

169. The Individual Defendants’ scienter is demonstrated by the fact that they intentionally changed the language in NVIDIA’s 2026 Proxy Statement pertaining specifically to whether the Company infringed copyright laws, removing language that had

been in NVIDIA's two prior proxy statements (for 2024 and 2025) to the effect that no such infringement had occurred. The following chart demonstrates the omission:

2024 Proxy	2025 Proxy	2026 Proxy
<i>Our trustworthy artificial intelligence, or AI, principles, which we share with customers and partners, reflect our core values and our Code of Conduct. We endeavor to deliver AI models that comply with privacy and data protection laws, perform safely and as intended, provide transparency about a model's design and limitations, minimize unwanted bias, and give equal opportunity to benefit from AI.</i>	<i>Our artificial intelligence, or AI, principles, which we share with customers and partners, reflect our core values and our Code of Conduct. We endeavor to deliver trustworthy AI models that comply with privacy and data protection laws, perform safely and as intended, provide transparency about a model's design and limitations, minimize unwanted bias, and give equal opportunity to benefit from AI.</i>	No such language.

170. The Individual Defendants' scienter is also demonstrated by their conduct in causing NVIDIA to repurchase billions of dollars of its own stock at inflated prices. During the time that NVIDIA was developing its LLMs, VLMs, and CVMs based on piracy datasets, the Board agreed to cause NVIDIA to repurchase its own common stock at prices that were artificially inflated by the materially false and misleading statements and omissions in NVIDIA's SEC filings.

171. In May 2022, August 2023, August 2024, and August 2025, the Board—while aware of the true extent of the Company's infringement on its AI platforms and datasets, and the false and misleading 2023-2025 Proxy Statements and 2023-2025 Annual Reports—authorized a share repurchase program of up to \$150 billion, initially through December 2023, then without expiration.

172. From June 2024 to January 25, 2026, NVIDIA repurchased \$13.258 billion of its own stock at significantly inflated prices through a combination of accelerated repurchase agreements and open-market purchases. All these repurchases were made at times when the Defendants knew, but had failed to disclose, the Company's copyright and BIPA infringement on its AI platforms and datasets. As a result, the Company paid artificially inflated prices to buy its own shares.

VI. DAMAGES TO NVIDIA

173. NVIDIA is significantly damaged by Defendants' misconduct. Due to Defendants' failure to take required steps to eliminate copyrighted works, YouTube videos, and commercialized voiceprints from copyrighted holders, content creators, and commercialized speakers from NVIDIA's AI datasets, NVIDIA will be subject to potentially massive liability from the copyright infringement and BIPA class actions. This includes not only the potential costs to resolve the litigation, but also legal fees, loss of goodwill, reputational damage, and lost customers, among other things.

174. NVIDIA was further damaged by Defendant Huang's unjust enrichment. Huang earned extensive compensation during the relevant period, as the Defendants' false statements and stock repurchases caused NVIDIA's stock to be inflated. Huang was awarded \$36.3 million in compensation for 2025, \$49.9 million in compensation for 2024, and \$34.2 million in compensation for 2023. Huang currently holds approximately 3.5% of NVIDIA's shares, a total of 851,983,603 shares, worth over \$173.9 billion.

175. As of July 12, 2025, in June and July 2025, Huang "unloaded roughly \$36.4 million worth of stock, or 225,000 shares," and his "wealth has grown ... about \$29 billion,

since the start of 2025 alone” with his net worth at \$143.7 billion.³⁰ As of October 1, 2025, Huang had sold more than \$1 billion worth of NVIDIA stock since June 2025.

176. “As of mid-2026, Huang’s net worth [was] estimated at around \$167 billion, well over twice his wealth at the end of 2023.”³¹ As of May 14, 2025, Huang’s net worth was \$200 billion, which did drop down to \$160 billion as of June 29, 2026, based on NVIDIA’s 20% stock drop.

177. All of the Officer Defendants received the following compensation for 2023-2025 (FY 2024-2026):

Name and Principal Position	Fiscal Year	Salary (\$)	Stock Awards (\$ (1))	Non-Equity Incentive Plan Compensation (\$ (2))	All Other Compensation (\$)	Total (\$ (3))
Jen-Hsun Huang	2026	1,497,627	24,800,511	6,000,000	4,045,691 (4)	36,343,830
<i>President and CEO</i>	2025	1,486,199	38,811,306	6,000,000	3,568,746	49,866,251
	2024	996,514	26,676,415	4,000,000	2,494,973	34, 167, 902
Colette M. Kress	2026	898,577	12,825,872	600,000	16,402 (5)	14,340,850
<i>EVP and CFO</i>	2025	893,739	19,849,891	600,000	18,902	21,362,532
	2024	896,863	11,756,027	600,000	13,902	13,266,792
Ajay K. Puri	2026	948,497	12,478,989	1,300,000	49,840 (5)	14,777,327

³⁰ Available at: <https://www.cnbc.com/2025/07/11/nvidia-stock-jensen-huang-stock-buffet.html>.

³¹ Available at: <https://www.thestreet.com/personalities/nvidia-founder-huang-net-worth>.

<i>EVP,</i>	2025	943,391	19,277,046	1,300,000	70,460	21,590,897
<i>Worldwide</i>						
<i>Field</i>						
<i>Operations</i>	2024	946,689	11,320,353	1,300,000	48,408	13,615,450
Debora						
Shoquist	2026	848,656	12,912,487	500,000	33,567 (5)	14,294,710
<i>EVP,</i>						
<i>Operations</i>	2025	844,087	17,838,832	500,000	34,984	19,217,903
	2024	847,037	9,687,599	500,000	24,229	11,058,865
Timothy S.						
Teter	2026	848,656	12,912,487	500,000	15,240 (5)	14,276,382
<i>EVP,</i>	2025	844,087	17,838,832	500,000	18,902	19,201,821
<i>General</i>						
<i>Counsel</i>						
<i>and</i>						
<i>Secretary</i>	2024	847,037	9,687,599	500,000	13,902	11,048,538

178. The Director Defendants received the following compensation for 2025:

Name	Fees Earned or Paid in Cash	Stock Awards (\$)	Total (\$)
Robert K. Burgess	63,750	278,809	342,559
Tench Coxe	85,000	278,809	363,809
John O. Dabiri	85,000	278,809	363,809
Persis S. Drell	85,000	278,809	363,809
Dawn Hudson	85,000	278,809	363,809
Harvey C. Jones	85,000	278,809	363,809
Melissa B. Lora	85,000	278,809	363,809
Stephen C. Neal	85,000	278,809	363,809
Ellen Ochoa	42,500	278,809	321,309
A. Brooke Seawell	85,000	278,809	363,809

Aarti Shah	85,000	278,809	363,809
Mark A. Stevens	85,000	278,809	363,809

179. The Director Defendants received the following compensation for 2024:

Name	Fees Earned or Paid in Cash	Stock Awards (\$)	Total (\$)
Robert K. Burgess	85,000	258,828	343,828
Tench Coxe	85,000	258,828	343,828
John O. Dabiri	85,000	258,828	343,828
Persis S. Drell	85,000	258,828	343,828
Dawn Hudson	85,000	258,828	343,828
Harvey C. Jones	85,000	258,828	343,828
Melissa B. Lora	85,000	258,828	343,828
Stephen C. Neal	85,000	278,809	343,828
Ellen Ochoa	32,550	439,659	472,209
A. Brooke Seawell	85,000	258,828	343,828
Aarti Shah	85,000	258,828	343,828
Mark A. Stevens	85,000	258,828	343,828

180. The Director Defendants received the following compensation for 2023:

Name	Fees Earned or Paid in Cash	Stock Awards (\$)	Total (\$)
Robert K. Burgess	85,000	274,268	359,268
Tench Coxe	85,000	274,268	359,268
John O. Dabiri	85,000	274,268	359,268
Persis S. Drell	85,000	274,268	359,268
Dawn Hudson	85,000	274,268	359,268
Harvey C. Jones	85,000	274,268	359,268

Melissa B. Lora	56,000	525,372	581,372
Stephen C. Neal	85,000	278,809	359,268
A. Brooke Seawell	85,000	274,268	359,268
Aarti Shah	85,000	274,268	359,268
Mark A. Stevens	85,000	274,268	359,268

181. NVIDIA was also damaged by spending nearly \$13.258 billion to repurchase roughly 87.1 million of its shares at market prices artificially inflated by the false and misleading statements and omissions in SEC filings approved by the Director Defendants, amounting to a massive waste of corporate assets.

VII. DERIVATIVE AND DEMAND FUTILITY ALLEGATIONS

182. NVIDIA is named as a Nominal Defendant solely in a derivative capacity. This is not a collusive action to confer jurisdiction on this Court that it would not otherwise have.

183. Plaintiff brings this action derivatively in the right and for the benefit of the Company to redress injuries suffered and to be suffered as a direct and proximate result of Defendants' wrongful conduct.

184. Plaintiff is a current stockholder of NVIDIA and has continuously owned NVIDIA common stock at all times relevant herein. Plaintiff will adequately and fairly represent the interests of the Company in enforcing and prosecuting its rights and has retained counsel competent and experienced in derivative litigation.

185. A pre-suit demand on the Board of NVIDIA is futile and, therefore, excused. At the time this action was commenced, the Board consisted of the ten Director Defendants: Huang, Coxe, Dabiri, Hudson, Jones, Lora, Neal, Seawell, Shaw, and Stevens. Because the current Board is comprised of an even number of directors, Plaintiff needs

only to allege demand futility as to half the Board members, or in this case five out of the ten directors.

186. Delaware law applies a director-by-director “three-part test as the universal test for assessing whether demand should be futile.” *United Food & Commercial Workers Union v. Zuckerberg*, 262 A.3d 1034, 1057-58 (Del. 2021). The three-part test asks with respect to each Director: “(i) whether the director received a material personal benefit from the alleged misconduct that is the subject of the litigation demand; (ii) whether the director would face a substantial likelihood of liability on any of the claims that are the subject of the litigation demand; and (iii) whether the director lacks independence from someone who received a material personal benefit from the alleged misconduct that is the subject of the litigation demand or who would face a substantial likelihood of liability on any of the claims that are the subject of the litigation demand.” *Id.*

A. Each Demand Director Faces a Substantial Likelihood of Personal Liability for Lack of Oversight

187. The second inquiry in the demand futility test under Delaware law is whether a director faces a substantial likelihood of liability on any of the claims that are the subject of the litigation demand. Here, each Director faces a substantial likelihood of liability for breaching his or her oversight duties in bad faith.

188. For a software company like NVIDIA, intellectual property rights, including copyright and biometric information protections, is a *mission critical issue* to the Company. If third parties could violate NVIDIA’s copyright and biometric information protections willy-nilly, NVIDIA’s software would have no value and the Company’s revenues would collapse. In a similar vein, if NVIDIA violates the copyright and biometric information rights of other companies or persons, it exposes itself to substantial damages. Even worse, if NVIDIA’s own software and datasets are developed in such a way that

violates intellectual property, copyright protection, and biometric information rights, the consequences are dire, as NVIDIA's software would be imbued with copyright and biometric information violations, posing a much more substantial threat, as the legal violations would threaten the Company's recurring revenues from its software, thereby not only exposing it to damages but also adversely impacting its revenues, profits, and growth rate. That is what happened here under the Director Defendants' watch.

189. All ten of the Directors—Huang, Coxe, Hudson, Jones, Lora, Neal, Seawell, Shah, and Stevens—have been on the Board of Directors at all relevant times, including during the time that NVIDIA committed the copyright and BIPA violations, and all Directors made the false and misleading statements and also caused NVIDIA to make the stock repurchases. Thus, all ten Directors are interested because they face a substantial likelihood of liability for each of the claims asserted herein.

190. The tenure of the current Board is set forth in this chart:

Director	Joined Board	Approx. Tenure
Jen-Hsun Huang	1993	33
Tench Coxe	1993	33
John Dabiri	2020	6
Dawn Hudson	2013	13
Harvey Jones	1993	33
Melissa Lora	2023	3
Stephen Neal	2019	7
Brooke Seawell	1997	29
Aarti Shah	2020	6
Mark Stevens	2008	18

191. The Directors also failed to act on numerous red flags related to NVIDIA's copyright and BIPA infringement. On March 11, 2024, a Law 360 article discussed how

“Writers Target Nvidia, Databricks in AI Copyright Litigation.” The article stated that NVIDIA was “accused of copyright infringements” by plaintiffs that “said they own registered copyrights in certain books that were included in the training dataset that Nvidia has admitted copying to train its NeMo Megatron models.”³² On January 21, 2026, the Directors were made aware that the Court agreed to permit these plaintiffs to amend their complaint—based on discovery in the *Nazemian* action—in an article entitled “NVIDIA must face direct infringement claim over ‘shadow library’ downloads.”³³ The article informed the Directors that plaintiffs had “newly discovered AI models allegedly trained on the illegally downloaded works” and produced “new claims of secondary infringement based on NVIDIA’s alleged acts encouraging its customers to download and use its training dataset.”

192. On December 1, 2025, the Directors were informed that “nVidia [has been] sued for allegedly scraping copyright-protected video from YouTube, in potential class action.”³⁴ The Piracy Monitor described the *Ted Entertainment* lawsuit, where “[t]he plaintiffs contend that nvidia[] violates the provisions of the Digital Millenium Copyright Act (DMCA) regarding anti-circumvention ... by bypassing technological protection measures that control access to copying of YouTube videos.” The Piracy Monitor news article explained:

³² Adam Lidgett, *Writers Target Nvidia, Databricks in AI Copyright Litigation*, LAW360, Mar. 11, 2024.

³³ Thomas Long, *NVIDIA must face direct infringement claim over ‘shadow library’ downloads*, VITALLAW, Jan. 21, 2026.

³⁴ *nVidia sued for allegedly scraping copyright-protected video from YouTube, in potential class action*, PIRACY MONITOR, Dec. 1, 2025.

What is the content in question?

All three plaintiffs contend that they have invested substantial time and money into producing their works and bringing awareness around it

- Ted Entertainment, Inc. is an independent content creator which owns and operates YouTube channels “h3h3 Productions” and “H3 Podcast Highlights,” with more than 5,800 original videos on YouTube, with a combined total of over 4,000,000,000 YouTube views and more than 2,600,000 subscribers.
- Plaintiff Matt Fisher is an individual creator and California resident who posts golf videos on YouTube, many of which are instructional. His Mr.ShortGame YouTube channel has over 500,000 subscribers and hundreds of millions of views.
- Golfholics is a California corporation with a golf channel that has 130,000 subscribers and has had millions of views.

Intentional infringement?

nVidia was said to use a video-downloading program combined with virtual machines that rotated IP addresses to avoid detection and blocking; which enabled the “mass extraction of, at times, eighty years of video content per day.”

Internal conversations detailed in the Complaint demonstrate that nVidia personnel were aware of the protections and actively worked toward circumventing them

193. On May 13, 2026, the Directors were informed of the *Rogers* action, “Illinois suits accuse Microsoft, Nvidia of stealing voiceprints for AI products.” In the Courthouse News Service article, Destiny Devooght explained that “Seven Illinois professionals claim Microsoft and Nvidia stole their unique voiceprints and used them to train commercial artificial intelligence voice synthesis models”:

The plaintiffs claim Nvidia ... obtained recordings of their voices, extracted the unique biometric signatures from the stolen voiceprints and used the data to build suites of commercial voice artificial intelligence products.

A voiceprint is the digital fingerprint of a human voice — a mathematical representation of its pitch, timbre, resonance, accent and distinctive speech patterns.

The companies used the voiceprints of highly accomplished professionals without consent or notification, according to the plaintiffs, knowingly violating Illinois' Biometric Information Privacy Act for the construction of products like ... Nvidia's Canary and Parakeet speech models.

The act requires written notice of the specific purpose and duration of voiceprint collection, along with a written release from the speaker.

Each suite of AI voice products utilized hundreds of thousands of hours of recordings from just as many speakers, including from each plaintiff, they say. Collection and training at that scale includes major compliance responsibilities, including identifying source speakers, providing notice and obtaining and storing countless releases.

194. All ten Directors approved each of the Company's 2023, 2024, 2025, and 2026 Proxy Statements. The Director all approved a major change in the language of the 2026 Proxy Statement that removed the prior, unequivocal statements in the prior years' proxy statements to the effect that NVIDIA was not infringing copyright and BIPA laws. This is strongly indicative of scienter. The following chart demonstrates the omission:

2024 Proxy	2025 Proxy	2026 Proxy
<i>Our trustworthy artificial intelligence, or AI, principles, which we share with customers and partners, reflect our core values and our Code of Conduct. We endeavor to deliver AI models that comply with privacy and data protection laws, perform safely and as intended, provide transparency about a model's design and limitations, minimize unwanted bias, and give equal opportunity to benefit from AI.</i>	<i>Our artificial intelligence, or AI, principles, which we share with customers and partners, reflect our core values and our Code of Conduct. We endeavor to deliver trustworthy AI models that comply with privacy and data protection laws, perform safely and as intended, provide transparency about a model's design and limitations, minimize unwanted bias, and give equal opportunity to benefit from AI.</i>	No such language.

195. The Director Defendants consciously ignored these red flags. For these reasons, the Director Defendants are incapable of making an independent and disinterested decision to institute and vigorously prosecute this action, including because they face a substantial likelihood of liability for their wrong acts, such that demand on the Board to institute this action is excused as futile.

B. Each Director Defendant Faces a Substantial Likelihood of Liability for Approving the Issuance of False and Misleading Statements

196. As explained above, the Individual Defendants caused NVIDIA to issue numerous false and misleading disclosures during the Relevant Period, misrepresenting that NVIDIA had not trained its AI with any copyrighted material and biometric information and concealing the extent to which NVIDIA had improperly used the copyrights and biometric information of other persons and companies.

197. The Director Defendants were ultimately responsible for the Company's public disclosures. Each Director Defendant approved one or more of the false and misleading statements alleged herein.

198. The Director Defendants together and individually, violated and breached their fiduciary duties of loyalty and acted in bad faith. The Director Defendants knowingly approved and/or permitted the wrongs alleged herein and caused the Company to engage in copyright infringement and BIPA violations, authorized and/or permitted the false statements to be disseminated directly to the public and made available and distributed to stockholders, authorized and/or permitted the issuance of various false and misleading statements, and are principal beneficiaries of the wrongdoing alleged herein, and thus, could not fairly and fully prosecute such a suit even if they instituted it.

199. The Director Defendants either knowingly or recklessly caused the Company to engage in copyright infringement and BIPA violation misconduct, and to issue

the materially false and misleading statements alleged herein. The Director Defendants knowingly approved and/or permitted a business model to train NVIDIA's AI products on a dataset that included copyrighted materials, YouTube videos, commercial voice models, and knew of the falsity of the misleading statements in SEC filings at the time they were made. The Director Defendants thus face a substantial likelihood of liability and are not disinterested.

200. As members of the Board charged with overseeing the Company's affairs, each of the Director Defendants had knowledge, or the fiduciary obligation to inform themselves, of information pertaining to the Company's operations and the material events giving rise to these claims. Specifically, as Board members of NVIDIA, the Director Defendants knew, or should have known, the material facts surrounding the Company's copyright infringement misconduct.

201. Moreover, the Director Defendants willfully ignored, or recklessly failed to inform themselves of, the obvious problems with the Company's internal controls, practices, and procedures, and failed to make a good faith effort to correct the problems or prevent their recurrence.

C. Defendant Huang is Interested and Received Improper Financial Benefits

202. Defendant Huang was the CEO of the Company at relevant times and is not disinterested or independent. Huang derived his sole employment income from NVIDIA and does not qualify as an independent director, as NVIDIA's proxy statements admit.

203. As alleged herein, Huang was awarded unjust compensation in the aggregate amount of \$120.40 million between 2023-25. Huang currently holds approximately 3.5% of NVIDIA's shares, 851,983,603 shares worth \$173.9 billion, based primarily on decades of his stock award compensation. His massive compensation was

unjust because he breached his fiduciary duties and issued false and misleading statements.

204. Further, Huang has led NVIDIA's establishment of the AI platforms, which subjects him substantially to the allegations of wrongdoing arising from many of the same events based on the same facts as alleged in this action, and the copyright infringement and BIPA violation actions. He is therefore interested in the claims asserted herein and demand is clearly futile as to Huang.

D. All Directors are Interested Because They Received Improper Financial Benefits Directly Related to Their Wrongdoing

205. Each of the Director Defendants is interested because they received improper financial benefits. To be clear, as particularly alleged herein, Plaintiff does not simply allege that the Directors Defendants received director fees. Instead, from 2023 to 2025 they received both cash fees and stock grants that were directly inflated due to the wrongdoing in which they directly participated.

206. Specifically, the Compensation Committee targets the equity compensation for the directors to be within the 25th to 75th percentile range of the peer group. The total shareholder return for NVIDIA, however, was well below that of the peer group identified in the proxy. In addition, the Director Defendants caused the Company to violate copyrights, BIPA, and IP, to issue false and misleading statements, and to repurchase billions of dollars of the Company's stock at inflated prices. This conduct caused NVIDIA's reported financial performance to be artificially inflated, thus directly inflating the financial metrics used by the Compensation Committee to set the Director Defendants' cash and stock compensation.

E. Demand is Futile as to the Audit Committee Directors

207. The Audit Committee Defendants (Coxe, Jones, Lora, Seawell, and Shah) are not disinterested or independent, and therefore, are incapable of considering a demand

because, pursuant to the Audit Committee Charter, were specifically charged with the responsibility to assist the Board in fulfilling its oversight responsibilities related to the adequacy and effectiveness of technology security policies, internal controls regarding information and technology security, cybersecurity, and privacy related areas, compliance with related legal, regulatory and ethical requirements, and public disclosure requirements. The Audit Committee Defendants breached their fiduciary duties to the Company by permitting and facilitating NVIDIA to conduct copyright and BIPA violations and failing to prevent, correct, or inform the Board of the issuance of material misstatements and omissions regarding the Company's copyright infringement and BIPA violation misconduct and the adequacy of the Company's internal controls as alleged above. The Audit Committee Defendants therefore cannot independently consider any demand to sue themselves for breaching their fiduciary duties to the Company and committing violations of the 1934 Act, as that would expose them to substantial liability and threaten their livelihood.

F. Demand is Futile for Additional Reasons

208. The Director Defendants, as members of the Board, were and are subject to the Company's Code of Conduct. The Code of Conduct goes well beyond the basic fiduciary duties required by applicable laws, rules, and regulations, requiring the Defendants to also adhere to NVIDIA's standards of business conduct. The Director Defendants violated the Code of Conduct because they knowingly or recklessly engaged in and participated in making and/or causing the Company to engage in copyright infringement and BIPA violation misconduct and make the materially false and misleading statements alleged herein. Because the Individual Defendants violated the Code of Conduct, they face a substantial likelihood of liability for breaching their fiduciary duties,

and therefore demand upon them is futile.

209. The Director Defendants' conduct was not the product of legitimate business judgment as it was based on bad faith and intentional, reckless, or disloyal misconduct. Thus, none of the directors can claim exculpation from their violations of duty pursuant to the Company's charter. As all of the directors face a substantial likelihood of liability, they are self-interested in the conduct challenged herein. They cannot be presumed to be capable of exercising independent and disinterested judgment about whether to pursue this action on behalf of the stockholders of the Company. Accordingly, demand is excused as being futile.

210. The acts complained of herein constitute violations of fiduciary duties owed by NVIDIA's officers and directors, and these acts are incapable of ratification.

211. The Director Defendants may also be protected against personal liability for their wrongful acts by directors' and officers' liability insurance if they caused the Company to purchase it for their protection with corporate funds—*i.e.*, monies belonging to NVIDIA. Any directors' and officers' liability insurance policy covering the Director Defendants will contain provisions that eliminate coverage for any action brought directly by the Company against its officers and directors, generally known as the "insured-versus-insured exclusion." As a result, if the Director Defendants were to sue themselves or certain officers of NVIDIA, there would be no directors' and officers' insurance protection. Accordingly, the Director Defendants cannot be expected to bring such a suit. On the other hand, if the suit is brought derivatively, as this action is brought, such insurance coverage will provide a basis for the Company to effectuate a recovery. Thus, demand on the Defendants is futile and, therefore, excused.

///

212. If there is no directors' and officers' liability insurance, then the directors will not cause NVIDIA to sue the Defendants named herein, since, if they did, they would face a large uninsured individual liability. Accordingly, demand is futile in that event as well.

213. Thus, there is no majority of disinterested directors among the Director Defendants and demand is excused as futile.

CAUSES OF ACTION

COUNT I

Against All Defendants for Breach of Fiduciary Duty

214. Plaintiff incorporates by reference and realleges each and every allegation contained above as though fully set forth herein.

215. Defendants owed and owe NVIDIA fiduciary obligations. By reason of their fiduciary relationships, the Defendants owed and owe NVIDIA the highest obligation of loyalty and care.

216. The Defendants and each of them, violated and breached their fiduciary duties.

217. The Officer Defendants either knew, were acting recklessly, and/or were grossly negligent in disregarding the unlawful activity of such substantial magnitude and duration. The Officer Defendants either knew, were reckless and/or were grossly negligent in not knowing that the Company and its statements concerning AI dataset were false and misleading. Accordingly, the Officer Defendants breached their duty of care and loyalty to the Company. The Officer Defendants are not entitled to the protection of the business judgment rule.

218. The Director Defendants breached their duty of loyalty by causing the Company to adopt and implement an unlawful business plan, pursuant to which NVIDIA

knowingly engaged in copyright infringement and BIPA violations in order to train its AI software. The strategy was carefully crafted to avoid spending the time and money necessary to obtain lawful licenses or permission from the copyright holders, and instead to simply pay damages later when NVIDIA was sued by the copyright holders. The Director Defendants knew or were reckless in not knowing that NVIDIA's AI dataset process was explicitly predicated upon widespread and blatant copyright infringement and BIPA violations. Accordingly, these defendants breached their duty of loyalty to the Company.

219. The Defendants also violated their fiduciary duties of loyalty and good faith by causing NVIDIA to repurchase millions of its own shares during the relevant time period, despite knowing that such prices were artificially inflated due to the false statements Defendants caused NVIDIA to make, as alleged more fully herein.

220. As a direct and proximate result of Defendants' breaches of their fiduciary obligations, NVIDIA has sustained significant damages, as alleged herein. As a result of the misconduct alleged herein, these defendants are liable to the Company.

221. Plaintiff, on behalf of NVIDIA, has no adequate remedy at law.

COUNT II

Against the Director Defendants for Violations of Section 14(a) of the Exchange Act and SEC Rule 14a-9

222. Plaintiff incorporates by reference and realleges each allegation contained above as though fully set forth herein, except to the extent those allegations plead fraud or intentional or knowing conduct by the Director Defendants. Plaintiff specifically disclaims any allegations of intentional or fraudulent misconduct with regard to this claim.

223. SEC Rule 14a-9, promulgated under Section 14(a) of the Exchange Act, prohibits solicitation through a proxy statement or other communication "containing any

statement which, at the time and in the light of the circumstances under which it is made, is false or misleading with respect to any material fact, or which omits to state any material fact necessary in order to make the statements therein not false or misleading or necessary to correct any statement in any earlier communication with respect to the solicitation of a proxy for the same meeting or subject matter which has become false or misleading.” The Director Defendants recklessly and/or negligently issued, caused to be issued, and participated in the issuance of materially misleading written statements to stockholders that were contained in the Proxy Statements. The Proxy Statements contained proposals to the Company’s stockholders urging them to re-elect the Director Defendants as members of the Board but misstated or failed to disclose the material information alleged herein.

224. By reasons of the conduct alleged herein, the Director Defendants violated Section 14(a) and Rule 14a-9. As a direct and proximate result of Defendants’ wrongful conduct, the Company misled or deceived its stockholders by making misleading statements that were an essential link in stockholders heeding the Company’s recommendations to re-elect the Board members up for election, approve the proposed executive compensation, and renew the contract of the outside auditor.

225. Plaintiff thereby seeks injunctive and equitable relief because the conduct of the Director Defendants is interfering with Plaintiff’s voting rights and choices at the annual meetings.

226. This action was timely commenced within three years of the date of the Proxy Statements and within one year from the time Plaintiff discovered or reasonably could have discovered the facts on which this claim is based.

///

///

COUNT III

Against the Director Defendants for Violations of Section 10(b) of the Exchange Act and SEC Rule 10b-5

227. Plaintiff incorporates by reference and realleges each and every allegation set forth above as though fully set forth herein.

228. During the Relevant Period, the Director Defendants disseminated or approved false or misleading statements and omissions about NVIDIA related to the Company's AI strategy, including representing that Defendants "endeavor to deliver trustworthy AI models that comply with privacy and protection laws" and "perform safely and as intended." Defendants made other false and misleading statements and omissions on similar topics, as alleged more particularly herein.

229. The Director Defendants knew or recklessly disregarded statements that were false or misleading and were intended to deceive, manipulate, or defraud. Those false or misleading statements and Defendants' course of conduct were designed to artificially inflate the price of the Company's common stock.

230. At the same time that the price of the Company's common stock was inflated due to the false or misleading statements, the Director Defendants caused the Company to repurchase billions of dollars of its own common stock at prices that were artificially inflated due to Defendants' false or misleading statements. From June 2024 to January 25, 2026, the Director Defendants caused NVIDIA to repurchase **\$13.258 billion** of its own stock at significantly inflated prices.

231. The Director Defendants violated Section 10(b) of the Exchange Act and SEC Rule 10b-5 in that they (a) employed devices, schemes, and artifices to defraud; (b) made untrue statements of material facts or omitted to state material facts necessary in order to make the statements made, in light of the circumstances under which they were

made, not misleading; and/or (c) engaged in acts, practices, and a course of business that operated as a fraud or deceit upon NVIDIA in connection with the Company's purchases of its stock during the Relevant Period.

232. The Director Defendants, individually and in concert, directly and indirectly, by the use of means or instrumentalities of interstate commerce or of the mails: (a) engaged and participated in a continuous course of conduct that operated as a fraud and deceit upon the Company; (b) made various false or misleading statements of material facts and omitted to state material facts necessary in order to make the statements made, in light of the circumstances under which they were made, not misleading; (c) made the above statements intentionally or with a severely reckless disregard for the truth; and (d) employed devices and artifices to defraud in connection with the purchase and sale of NVIDIA stock, which were intended to, and did, (i) deceive NVIDIA regarding, among other things, the Company's lack of internal controls to monitor, detect and prevent illegal copyright infringement; and (ii) artificially inflate and maintain the market price of NVIDIA stock.

233. The Director Defendants were among the senior management and the directors of the Company, and were therefore directly responsible for, and are liable for, all materially false or misleading statements made during the Relevant Period, as alleged above.

234. As described above, Defendants acted with scienter in that they acted either with intent to deceive, manipulate, or defraud, or with severe recklessness. The misstatements and omissions of material facts set forth herein were either known to Defendants or were so obvious that Defendants should have been aware of them. Throughout the Relevant Period, Defendants also had a duty to disclose new information

that came to their attention and rendered their prior statements materially false or misleading.

235. NVIDIA has and will suffer damages in that it paid artificially inflated prices for NVIDIA common stock purchased as part of the repurchase program and suffered losses when the previously undisclosed facts relating to the Company's copyright infringements and BIPA violations were disclosed, though the full extent of such wrongdoing has not yet been revealed. NVIDIA would not have purchased these securities at the prices it paid, or at all, but for the artificial inflation in the Company's stock price caused by the false or misleading statements.

236. By reason of such conduct, Defendants are liable to the Company pursuant to Section 10(b) of the Exchange Act and SEC Rule 10b-5.

237. This action was timely commenced within two years from the time Plaintiff discovered or reasonably could have discovered the facts on which this claim is based and within five years of the violation.

COUNT IV
Violation of § 20A of The Securities Exchange Act of 1934
(Against Defendant Huang)

238. Plaintiff incorporates by reference and realleges each and every allegation contained above, as though fully set forth in this paragraph.

239. While NVIDIA's securities traded at artificially inflated and distorted prices, Defendant Huang personally profited by selling shares of NVIDIA common stock on the dates alleged *supra*—while he was in possession of adverse, material non-public information about the Company. As of July 12, 2025, Huang sold 225,000 shares for proceeds of \$36.4 million, and from July 15, 2025 to October 29, 2025, Huang sold 5,250,000 shares for proceeds over \$1 billion. As of the date of the filing of this Complaint,

Huang has received proceeds of approximately \$2.9 billion for NVIDIA stock sales.

240. Contemporaneously, NVIDIA repurchased its own shares during this same period, repurchasing \$13.258 billion of its own stock at significantly inflated prices.

241. The exact dates of each repurchase of NVIDIA stock by the Company is not known, and discovery is necessary to learn all the exact dates because NVIDIA does not disclose the details of each stock repurchase, and instead only publishes summary data. The data publicly disclosed by NVIDIA is alleged herein. However, based on the huge volume of shares involved in NVIDIA's stock repurchase plan, and the fact that the repurchase program was ongoing during the entire Relevant Time Period, Plaintiff believes, and on that basis alleges, that NVIDIA repurchased its shares within one to three days of each date that the Defendant sold stock.

242. NVIDIA suffered damages because: (a) in reliance on the integrity of the market, NVIDIA paid artificially inflated prices as a result of Defendants' violations of §§ 10(b) and 20(a) of the 1934 Act as alleged herein; and (b) NVIDIA would not have repurchased its shares at the same prices if it had been aware that the market for NVIDIA stock had been artificially inflated by the false and misleading statements and omissions alleged herein.

243. By reason of the foregoing, Defendant Huang violated § 20A of the 1934 Act and is liable to NVIDIA for the substantial damages the Company sustained in connection with its purchases of NVIDIA common stock during the Relevant Period.

VIII. PRAYER FOR RELIEF

WHEREFORE, Plaintiff prays for judgment in favor of Plaintiff and the Company and against all Individual Defendants as follows:

A. Declaring and determining that Plaintiff may maintain this action on behalf of NVIDIA and that Plaintiff is an adequate representative of the Company;

B. Declaring and determining that the Defendants breached their fiduciary duties to NVIDIA;

C. Declaring and determining that the Director Defendants violated Section 14(a) of the Exchange Act and SEC Rule 14a-9 by reason of the acts and omissions alleged herein;

D. Declaring and determining that the Director Defendants violated Section 10(b) of the Exchange Act and SEC Rule 10b-5 by reason of the acts and omissions alleged herein;

E. Declaring and determining that Huang violated Section 20(A) of the Exchange Act by reason of the acts and omissions alleged herein;

F. Determining and awarding to NVIDIA the damages sustained by it as a result of the violations set forth above from each of the Defendants, jointly and severally, together with pre-judgment and post-judgment interest thereon;

G. Directing NVIDIA and the Defendants to take all necessary actions to reform and improve its corporate governance and internal procedures to comply with applicable laws and to protect NVIDIA and its stockholders from a repeat of the damaging events described herein, including, but not limited to, actions as may be necessary to ensure proper corporate governance policies to strengthen the internal audit and control functions and ensure the establishment of effective oversight of compliance with copyright laws and

all other applicable laws, rules, and regulations.

H. Awarding NVIDIA restitution from Defendants, and disgorgement of any benefits wrongly obtained by Defendants;

I. Awarding Plaintiff the costs and disbursements of this action, including reasonable attorneys' and experts' fees, costs, and expenses; and

J. Granting such other and further relief as the Court may deem just and proper.

JURY TRIAL DEMAND

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

Dated: July 31, 2026

Respectfully submitted,

s/ David T. Wissbroecker

David T. Wissbroecker

BOTTINI & BOTTINI, INC.
Francis A. Bottini, Jr.
David T. Wissbroecker
7817 Ivanhoe Avenue, Suite 102
La Jolla, CA 92037
Telephone: (858) 914-2001
Facsimile: (858) 914-2002
fbottini@bottinilaw.com
dwissbroecker@bottinilaw.com

VERIFICATION

I, Jessica Berliner, verify that I am a shareholder of record of NVIDIA Corporation. I have reviewed the allegations in this Verified Stockholder Derivative Complaint. As to those allegations of which I have personal knowledge, I believe them to be true; as to those allegations of which I lack personal knowledge, I rely upon my counsel and counsel's investigation, and believe them to be true. Having received a copy of the complaint and reviewed it with counsel, I authorize its filing.

I declare under penalty of perjury under the laws of the United States of America that the foregoing is true and correct. Executed on this 29th day of July, 2026.

DocuSigned by:



D950EB38AA5C495

Jessica Berliner