



This report may not be used, reproduced, or distributed by any person, in whole or in part or in any form or manner, including creating any summaries thereof, without Glass Lewis' prior express written consent.

NASDAQ: **MSFT**

ISIN: **US5949181045**

MEETING DATE: 10 DECEMBER 2024

RECORD DATE: 30 SEPTEMBER 2024

PUBLISH DATE: 19 NOVEMBER 2024

INDEX MEMBERSHIP: DOW JONES COMPOSITE AVERAGE;
RUSSELL 3000; NASDAQ COMPOSITE;
S&P GLOBAL 100; DOW JONES
INDUSTRIAL AVERAGE; NASDAQ-100;
RUSSELL 1000; S&P 100; S&P 500

COMPANY DESCRIPTION

Microsoft Corporation develops and supports software, services, devices and solutions worldwide.

SECTOR: INFORMATION TECHNOLOGY

INDUSTRY: SOFTWARE

COUNTRY OF TRADE: UNITED STATES

COUNTRY OF INCORPORATION: UNITED STATES

HEADQUARTERS: WASHINGTON

VOTING IMPEDIMENT: NONE

OWNERSHIP	COMPANY PROFILE	ESG PROFILE	COMPENSATION	COMPENSATION ANALYSIS	COMPANY UPDATES
PEER COMPARISON	VOTE RESULTS	COMPANY FEEDBACK	APPENDIX	SUSTAINALYTICS ESG	ESG BOOK PROFILE
BITSIGHT CYBER SECURITY					

 2024 ANNUAL MEETING

PROPOSAL	ISSUE	BOARD	GLASS LEWIS	CONCERNS
1.00	Election of Directors	FOR	SPLIT	
1.01	Elect Reid G. Hoffman	FOR	FOR	
1.02	Elect Hugh F. Johnston	FOR	AGAINST	• Overboarded
1.03	Elect Teri L. List	FOR	FOR	
1.04	Elect Catherine MacGregor	FOR	FOR	
1.05	Elect Mark Mason	FOR	FOR	
1.06	Elect Satya Nadella	FOR	FOR	
1.07	Elect Sandra E. Peterson	FOR	FOR	
1.08	Elect Penny S. Pritzker	FOR	FOR	
1.09	Elect Carlos A. Rodriguez	FOR	FOR	
1.10	Elect Charles W. Scharf	FOR	FOR	
1.11	Elect John W. Stanton	FOR	FOR	
1.12	Elect Emma N. Walmsley	FOR	FOR	
2.00	Advisory Vote on Executive Compensation	FOR	FOR	
3.00	Ratification of Auditor	FOR	FOR	

4.00	Shareholder Proposal Regarding Risks of Developing Military Weapons	AGAINST	FOR	Additional disclosure could help shareholders understand financial and reputational risks from work with the military
5.00	Shareholder Proposal Regarding Assessment of Investments in Bitcoin	AGAINST	AGAINST	
6.00	Shareholder Proposal Regarding Report on Siting in Countries of Significant Human Rights Concern	AGAINST	AGAINST	
7.00	Shareholder Proposal Regarding Report on Risks of Providing AI to Facilitate New Oil and Gas Development and Production	AGAINST	AGAINST	
8.00	Shareholder Proposal Regarding Report on AI Misinformation and Disinformation	AGAINST	FOR	Information concerning exposure to risks related to misinformation and disinformation could be decision-useful for shareholders
9.00	Shareholder Proposal Regarding Report on Risks of AI Data Sourcing	AGAINST	FOR	Additional disclosure will better allow shareholders to understand the Company's management of AI-related risks

POTENTIAL CONFLICTS

As of October 2021, U.S. and Canadian companies are eligible to purchase and receive Equity Plan Advisory services from Glass Lewis Corporate, LLC ("GLC"), a Glass Lewis affiliated company. More information, including whether the company that is the subject of this report used GLC's services with respect to any equity plan discussed in this report, is available to Glass Lewis' institutional clients on Viewpoint or by contacting compliance@glasslewis.com. Glass Lewis maintains a strict separation between GLC and its research analysts. GLC and its personnel did not participate in any way in the preparation of this report.

ENGAGEMENT ACTIVITIES

Glass Lewis held the following engagement meetings within the past year:

ENGAGED WITH	MEETING DATE	ORGANIZER	TYPE OF MEETING	TOPICS DISCUSSED
Issuer	08 October 2024	Issuer	Teleconference/Web-Meeting	Board Composition and Performance, Executive Pay, Shareholder Proposal
Shareholder Proponent	10 October 2024	Shareholder Proposal Proponent	Teleconference/Web-Meeting	Shareholder Proposal

For further information regarding our engagement policy, please visit <http://www.glasslewis.com/engagement-policy/>.

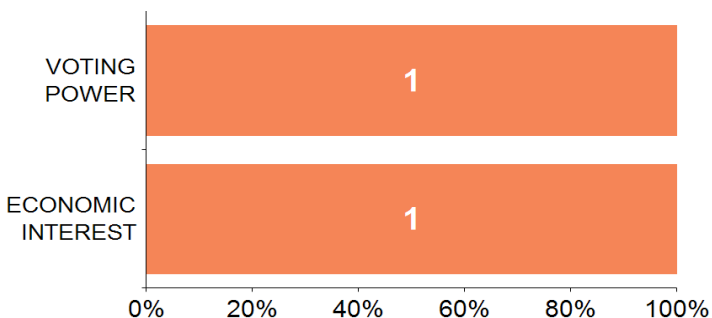
ISSUER DATA REPORT: Microsoft Corporation registered to participate in Glass Lewis' Issuer Data Report program (IDR) for this meeting. The IDR program enables companies to preview the key data points used by Glass Lewis' research team, and address any factual errors with Glass Lewis prior to the publication of the Proxy Paper to Glass Lewis' clients. No voting recommendations or analyses are provided as part of the IDR. For more information on the IDR program, please visit <https://www.glasslewis.com/issuer-data-report/>

SHARE OWNERSHIP PROFILE

SHARE BREAKDOWN

	1
SHARE CLASS	Common Stock
SHARES OUTSTANDING	7,434.9 M
VOTES PER SHARE	1
INSIDE OWNERSHIP	0.00%
STRATEGIC OWNERS**	0.10%
FREE FLOAT	99.90%

SOURCE CAPITAL IQ AND GLASS LEWIS. AS OF 19-NOV-2024



TOP 20 SHAREHOLDERS

	HOLDER	OWNED*	COUNTRY	INVESTOR TYPE
1.	The Vanguard Group, Inc.	9.06%	United States	Traditional Investment Manager
2.	BlackRock, Inc.	7.45%	United States	Traditional Investment Manager
3.	State Street Global Advisors, Inc.	3.89%	United States	Traditional Investment Manager
4.	Capital Research and Management Company	3.02%	United States	Traditional Investment Manager
5.	FMR LLC	2.57%	United States	Traditional Investment Manager
6.	Geode Capital Management, LLC	2.22%	United States	Traditional Investment Manager
7.	T. Rowe Price Group, Inc.	1.98%	United States	Traditional Investment Manager
8.	JP Morgan Asset Management	1.31%	United States	Traditional Investment Manager
9.	Norges Bank Investment Management	1.31%	Norway	Government Pension Plan Sponsor
10.	UBS Asset Management AG	1.17%	Switzerland	Traditional Investment Manager
11.	Northern Trust Global Investments	0.97%	United Kingdom	Traditional Investment Manager
12.	Legal & General Investment Management Limited	0.88%	United Kingdom	Traditional Investment Manager
13.	Morgan Stanley, Investment Banking and Brokerage Investments	0.86%	United Kingdom	Bank/Investment Bank
14.	BNY Asset Management	0.76%	United States	Traditional Investment Manager
15.	Wellington Management Group LLP	0.73%	United States	Traditional Investment Manager
16.	Teachers Insurance and Annuity Association-College Retirement Equities Fund	0.72%	United States	Traditional Investment Manager
17.	Eaton Vance Management	0.67%	United States	Traditional Investment Manager
18.	Charles Schwab Investment Management, Inc.	0.64%	United States	Traditional Investment Manager
19.	AllianceBernstein L.P.	0.55%	United States	Traditional Investment Manager
20.	Goldman Sachs Asset Management, L.P.	0.51%	United States	Traditional Investment Manager

*COMMON STOCK EQUIVALENTS (AGGREGATE ECONOMIC INTEREST) SOURCE: CAPITAL IQ. AS OF 19-NOV-2024

**CAPITAL IQ DEFINES STRATEGIC SHAREHOLDER AS A PUBLIC OR PRIVATE CORPORATION, INDIVIDUAL/INSIDER, COMPANY CONTROLLED FOUNDATION, ESOP OR STATE OWNED SHARES OR ANY HEDGE FUND MANAGERS, VC/PE FIRMS OR SOVEREIGN WEALTH FUNDS WITH A STAKE GREATER THAN 5%.

SHAREHOLDER RIGHTS

	MARKET THRESHOLD	COMPANY THRESHOLD ¹
VOTING POWER REQUIRED TO CALL A SPECIAL MEETING	N/A	15.00%
VOTING POWER REQUIRED TO ADD AGENDA ITEM	\$2,000 ²	\$2,000 ²
VOTING POWER REQUIRED TO APPROVE A WRITTEN CONSENT	N/A	100.00%

¹N/A INDICATES THAT THE COMPANY DOES NOT PROVIDE THE CORRESPONDING SHAREHOLDER RIGHT.

²UNLESS GRANDFATHERED, SHAREHOLDERS MUST OWN SHARES WITH MARKET VALUE OF AT LEAST \$2,000 FOR THREE YEARS. ALTERNATIVELY, SHAREHOLDERS MUST OWN SHARES WITH MARKET VALUE OF AT LEAST \$15,000 FOR TWO YEARS; OR SHARES WITH MARKET VALUE OF \$25,000 FOR AT LEAST ONE YEAR.

COMPANY PROFILE

FINANCIALS		1 YR TSR	3 YR TSR AVG.	5 YR TSR AVG.
	MSFT	32.3%	19.2%	28.5%
	S&P 500	24.6%	10.0%	15.0%
	Peers*	28.9%	15.0%	24.1%
	Market Capitalization (MM \$)	3,321,869		
	Enterprise Value (MM \$)	3,401,406		
	Revenues (MM \$)	245,122		

ANNUALIZED SHAREHOLDER RETURNS: *PEERS ARE BASED ON THE INDUSTRY SEGMENTATION OF THE GLOBAL INDUSTRIAL CLASSIFICATION SYSTEM (GICS). FIGURES AS OF 30-JUN-2024. SOURCE: CAPITAL IQ

EXECUTIVE COMPENSATION	Total CEO Compensation \$79,106,183			
	1-Year Change in CEO Pay	63%	CEO to Median Employee Pay Ratio	408:1
	Say on Pay Frequency	1 Year	Compensation Grade 2024	C
	Glass Lewis Structure Rating	Fair	Glass Lewis Disclosure Rating	Fair
	Single Trigger CIC Vesting	No	Excise Tax Gross-Ups	No
	NEO Ownership Guidelines	Yes	Overhang of Incentive Plans	2.93%

CORPORATE GOVERNANCE	Election Method	Majority	CEO Start Date	February 2014
	Controlled Company	No	Proxy Access	Yes
	Multi-Class Voting	No	Virtual-Only Meeting	Yes
	Staggered Board	No	Average NED Tenure	6 years
	Combined Chair/CEO	Yes	Gender Diversity on Board	41.7%
	Individual Director Skills Matrix Disclosed	Yes	Company-Reported Racial/Ethnic Diversity on Board	25.0%
	Supermajority* to Amend Bylaws and/or Charter	No	Age-Based Director Retirement Policy/Guideline	Yes; 75
	Numerical Director Commitments Policy	Yes		

*Supermajority defined as at least two-thirds of shares outstanding

ANTI-TAKEOVER	Poison Pill	No
	Approved by Shareholders/Expiration Date	N/A; N/A

AUDITORS	Auditor: DELOITTE & TOUCHE	Tenure: 41 Years
	Material Weakness(es) Outstanding	No
	Restatement(s) in Past 12 Months	No

SASB MATERIALITY	Primary SASB Industry: Software & IT Services		
	Financially Material Topics:		
	- Environmental Footprint of Hardware Infrastructure - Recruiting & Managing a Global, Diverse & Skilled Workforce - Managing Systemic Risks from Technology Disruptions - Data Privacy & Freedom of Expression - Data Security - Intellectual Property Protection & Competitive Behavior		
	Company Reports to SASB/Extent of Disclosure: Yes; Full Standard		

CURRENT AS OF NOV 19, 2024

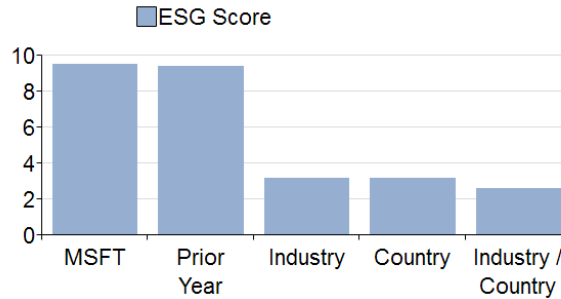
GLASS LEWIS ESG PROFILE

GLASS LEWIS ESG SCORE: 9.5 / 10

ESG SCORE SUMMARY

Board Accountability Score:	10.0 / 10	ESG Transparency Score:	8.8 / 10	Targets and Alignment Score:	10.0 / 10
Climate Risk Mitigation Score:	N/A	Biodiversity Score:	N/A		

SCORE BREAKDOWN



PRIOR YEAR ESG SCORE*	9.418
CHANGE IN ESG SCORE	0.13
INDUSTRY	3.2 (6.37)
COUNTRY	3.2 (6.35)
INDUSTRY / COUNTRY	2.6 (6.93)

*As of our Proxy Paper for the Annual Meeting on 07-Dec-23

BOARD ACCOUNTABILITY (10.0 / 10)

Average NED Tenure	6 years	Percent Gender Diversity	42%
Director Independence	92%	Board Oversight of ESG	Yes
Board Oversight of Cyber	Yes	Board Oversight of Human Capital	Yes
Compensation Linked to E&S Metrics	Yes	Lowest Support for Directors in Prior Year	91.1%
Prior Year Say on Pay Support	93.3%	Annual Director Elections	Yes
Inequitable Voting Rights	No	Pay Ratio	408:1
Diversity Disclosure Assessment	Exemplary	Failure to Respond to Shareholder Proposal	No

ESG TRANSPARENCY (8.8 / 10)

Comprehensive Sustainability Reporting	Yes	GRI-Indicated Report	No
Reporting Assurance	Yes	Reporting Aligns with TCFD	Yes
Discloses Scope 1 & 2 Emissions	Yes	Discloses Scope 3 Emissions	Yes
Reports to SASB	Yes	Extent of SASB Reporting	Full Standard
Discloses EEO-1 Report	Yes	CPA-Zicklin Score	88.6

ESG TARGETS AND ALIGNMENT (10.0 / 10)

Has Scope 1 and/or 2 GHG Reduction Targets	Yes	Has Scope 3 GHG Reduction Targets	Yes
Has Net Zero GHG Target	Yes	Reduction Target Certified by SBTi	Yes
SBTi Near-Term Target	1.5 degrees	SBTi Long-Term Target	N/A
SBTi Net Zero Target	Commitment Removed	UNGC Participant or Signatory	Yes
Has Human Rights Policy	Yes	Human Rights Policy Aligns with UDHR or ILO	Yes
Has Biodiversity Policy	Yes		

© 2024 Glass, Lewis & Co., and/or its affiliates. All Rights Reserved. The use of, or reference to, any data point, metric, or score collected, issued, or otherwise provided by a third-party company or organization (each, a "Third Party"), or a reference to such Third Party itself, in no way represents or implies an endorsement, recommendation, or sponsorship by such Third Party of the ESG Profile, the ESG Score, any methodology used by Glass Lewis, Glass Lewis itself, or any other Glass Lewis products or services. For further details about our methodology and data included in this page please refer to our [methodology documentation here](#).

PAY-FOR-PERFORMANCE

Microsoft's executive compensation received a **C** grade in our proprietary pay-for-performance model. The Company paid more compensation to its named executive officers than the median compensation for a group of companies selected based on Glass Lewis' peer group methodology and Diligent Intel's company data. The CEO was paid more than the median CEO compensation of these peer companies. Overall, the Company paid more than its peers and performed better than its peers.

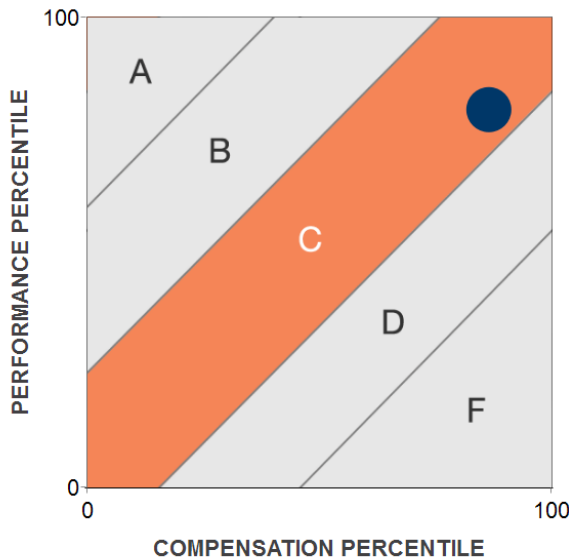
HISTORICAL COMPENSATION GRADE

FY 2024:	C
FY 2023:	C
FY 2022:	C

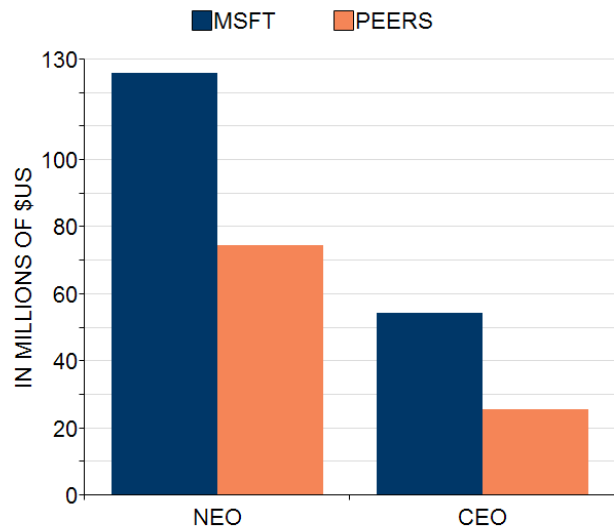
FY 2024 CEO COMPENSATION

SALARY:	\$2,500,000
GDFV EQUITY:	\$55,874,717
NEIP/OTHER:	\$5,369,791
TOTAL:	\$63,744,508

FY 2024 PAY-FOR-PERFORMANCE GRADE



3-YEAR WEIGHTED AVERAGE COMPENSATION

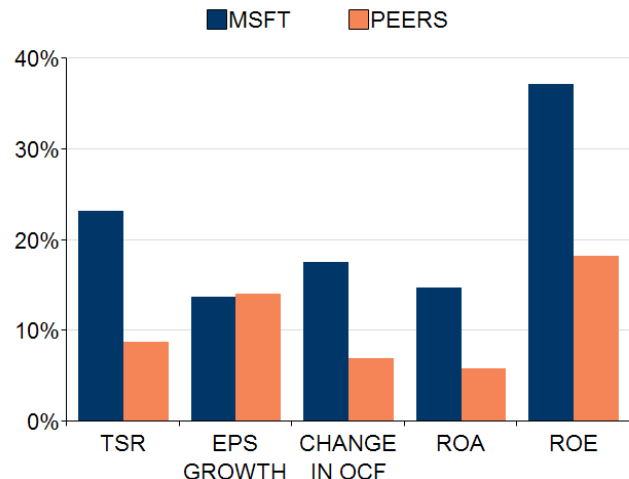


GLASS LEWIS PEERS VS PEERS DISCLOSED BY COMPANY

GLASS LEWIS	MSFT
Alphabet Inc.*	Netflix, Inc.
Apple Inc.*	Salesforce, Inc.
Amazon.com, Inc.*	QUALCOMM Incorporated
Meta Platforms, Inc.*	The Procter & Gamble Company
Johnson & Johnson*	Pfizer Inc.
Verizon Communications Inc.*	Oracle Corporation
AT&T Inc.*	Merck & Co., Inc.
The Walt Disney Company*	Intel Corporation
International Business Machines Corporation*	Cisco Systems, Inc.
Walmart Inc.	Adobe Inc.
Comcast Corporation*	Accenture PLC
The Boeing Company	
UnitedHealth Group Incorporated	
NVIDIA Corporation*	
Tesla, Inc.*	

***ALSO DISCLOSED BY MSFT**

SHAREHOLDER WEALTH AND BUSINESS PERFORMANCE



Analysis for the year ended 6/30/2024. Performance measures, except ROA and ROE, are based on the weighted average of annualized one-, two- and three-year data. Compensation figures are weighted average three-year data calculated by Glass Lewis. Data for Glass Lewis' pay-for-performance tests are sourced from Diligent Compensation & Governance Intel and company filings, including proxy statements, annual reports, and other forms for pay. Performance and TSR data are sourced from Capital IQ and publicly filed annual reports. For Canadian peers, equity awards are normalized using the grant date exchange rate and cash compensation data is normalized using the fiscal year-end exchange rate. The performance metrics used in the analysis are selected by Glass Lewis and standardized across companies by industry. These metrics may differ from the key metrics disclosed by individual companies to meet SEC pay-versus-performance rules.

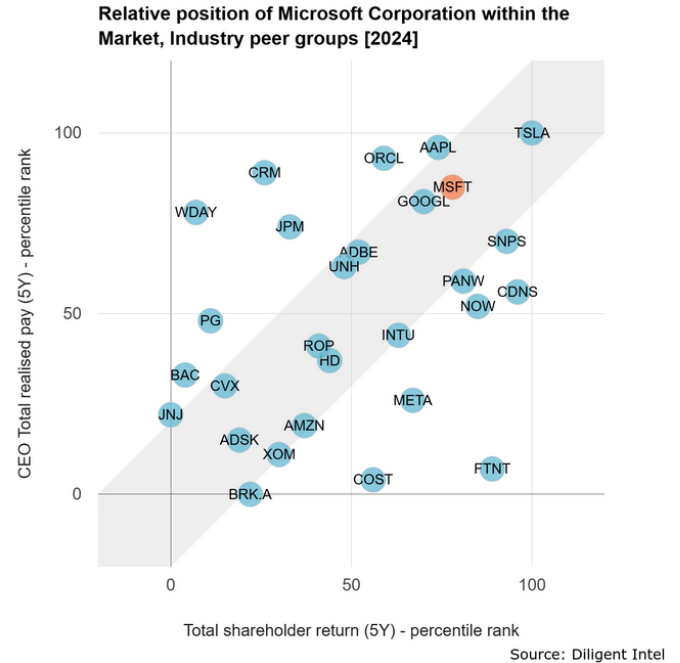
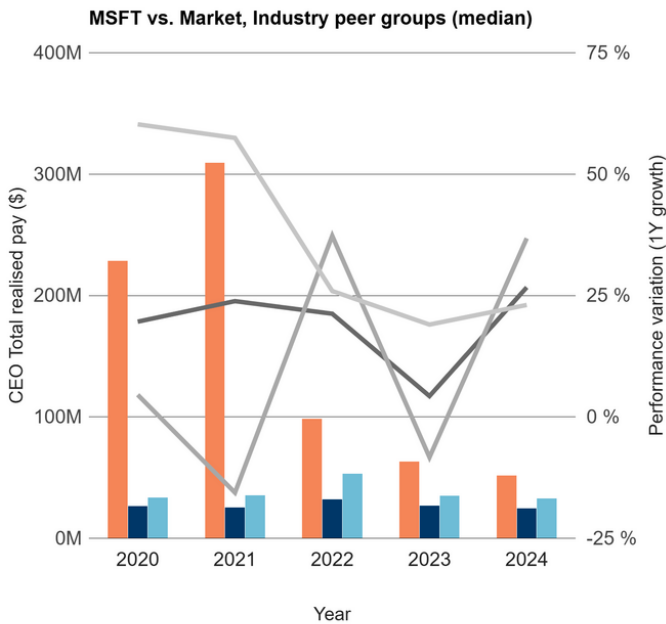
Glass Lewis peers are based on Glass Lewis' proprietary peer methodology, which considers both country-based and sector-based peers, along with each company's disclosed peers, and are updated in February and August. Peer data is based on publicly available information, as well as information provided to Glass Lewis during the open submission periods. The "Peers Disclosed by Company" data is based on public information in proxy statements and on companies' submissions. Glass Lewis may

exclude certain peers from the Pay for Performance analysis based on factors such as trading status and/or data availability.

For details on the Pay-for-Performance analysis and peer group methodology, please refer to Glass Lewis' [Pay-for-Performance Methodology & FAQ](#).

The intellectual property rights to the Diligent Compensation & Governance Intel data are vested exclusively in Diligent Corporation. Diligent Corporation and its affiliates and suppliers do not make any representation or warranty, express or implied, of any nature, and do not accept any responsibility or liability of any kind, including with respect to the accuracy, completeness or suitability for any purpose of the information contained herein arising from or relating to the use of the Diligent Compensation & Governance Intel data in connection with this Proxy Paper in any manner whatsoever.

COMPENSATION ANALYSIS



Total realised pay (MSFT)	Total realised pay (Market)	Total realised pay (Industry)	EBITDA (MSFT)	EBITDA (Market)	EBITDA (Industry)

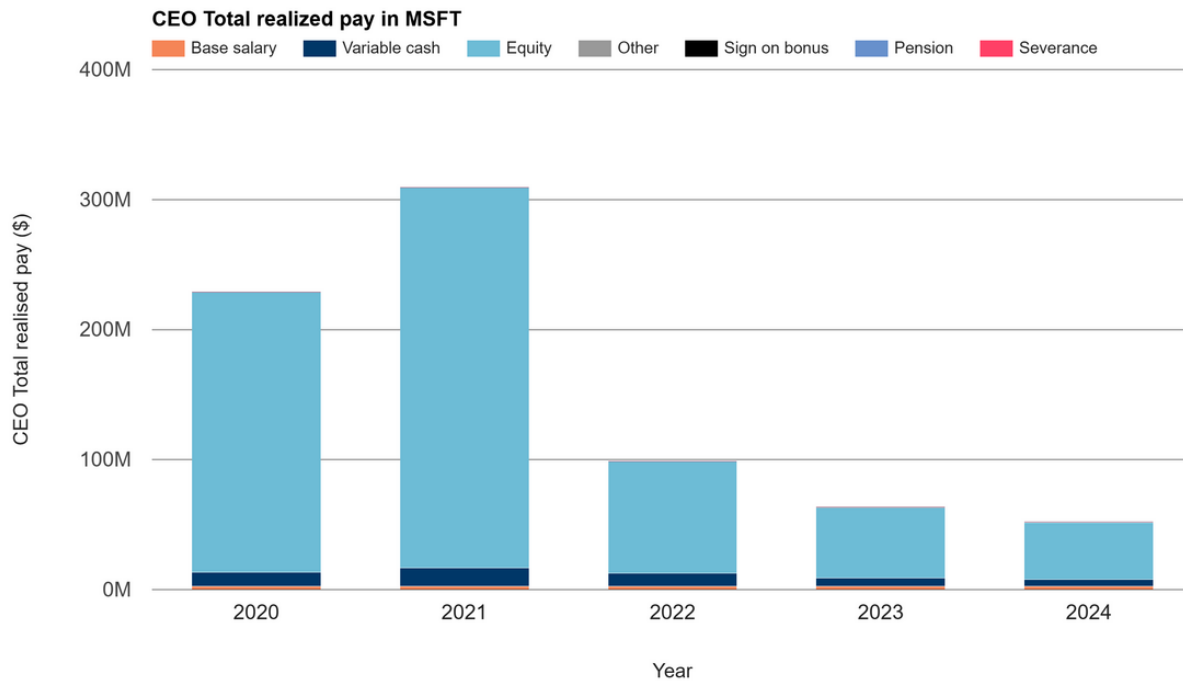
* All financial metrics are plotted at fiscal year growth rates in the graphs above. Absolute values are found in the tables below.

Year	Total realised pay (\$)*			EBITDA (\$)*			ROA			ROIC		
	MSFT	Market (Median)	Industry (Median)	MSFT	Market (Median)	Industry (Median)	MSFT	Market (Median)	Industry (Median)	MSFT	Market (Median)	Industry (Median)
2024	51.9	24.9	33.0	129,433.0	42,177.0	1,349.9	14.8%	8.0%	4.0%	21.0%	13.6%	5.9%
2023	63.4	27.1	35.1	102,175.0	30,835.0	1,097.0	14.3%	10.0%	3.8%	20.9%	14.6%	4.5%
2022	98.6	32.4	53.5	97,983.0	33,632.0	921.8	14.9%	7.4%	5.7%	22.2%	12.9%	8.3%
2021	309.4	25.5	35.6	80,816.0	24,496.0	732.0	13.8%	6.7%	5.7%	20.6%	12.9%	8.4%
2020	228.8	26.6	33.9	65,259.0	29,024.0	464.9	11.3%	7.5%	4.3%	17.0%	12.6%	6.0%

* Values provided in millions.

List of companies

Market peer group	Alphabet Inc. (GOOGL), Amazon.com, Inc. (AMZN), Apple Inc. (AAPL), Bank of America Corporation (BAC), Berkshire Hathaway Inc. (BRK.A), Chevron Corporation (CVX), Costco Wholesale Corporation (COST), Exxon Mobil Corporation (XOM), Johnson & Johnson (JNJ), JPMorgan Chase & Co. (JPM), Meta Platforms, Inc. (META), Tesla, Inc. (TSLA), The Home Depot, Inc. (HD), The Procter & Gamble Company (PG), UnitedHealth Group Incorporated (UNH)
Industry peer group	Adobe Inc. (ADBE), Atlassian Corporation (TEAM), Autodesk, Inc. (ADSK), Cadence Design Systems, Inc. (CDNS), CrowdStrike Holdings, Inc. (CRWD), Fortinet, Inc. (FTNT), Intuit Inc. (INTU), Oracle Corporation (ORCL), Palo Alto Networks, Inc. (PANW), Roper Technologies, Inc. (ROP), Salesforce, Inc. (CRM), ServiceNow, Inc. (NOW), Synopsys, Inc. (SNPS), Workday, Inc. (WDAY), Zoom Video Communications, Inc. (ZM)



Source: Diligent Intel

Year	Total realised pay (\$)	Base salary (\$)	Variable cash (\$)	Equity (\$)	Other (\$)	Sign on bonus (\$)	Pension (\$)	Severance (\$)
2024	51,881,567	2,500,000	5,200,000	44,011,776	169,791	0	0	0
2023	63,398,991	2,500,000	6,414,750	54,122,591	361,650	0	0	0
2022	98,591,001	2,500,000	10,066,500	85,914,251	110,250	0	0	0
2021	309,368,186	2,500,000	14,212,500	292,545,936	109,750	0	0	0
2020	228,755,029	2,500,000	10,992,000	215,151,849	111,180	0	0	0

For further information on the peers and methodology, or to submit feedback, please see our [FAQs](#).

The Compensation Analysis is based on Glass Lewis' proprietary methodology using Diligent Intel proprietary platform. The intellectual property rights to the platform are vested exclusively in Diligent Compensation & Governance Intel, the brand under which Diligent Corporation operates and provides these services. Compensation figures are standardized and calculated by Diligent Intel based on information disclosed by the Company and its peers in their disclosures and proxy materials. For realizable pay reported for European and Australian companies, equity awards are normalized using the vesting date share price or when not disclosed by the Company using the year end share price. For U.S. and Canadian companies, realized pay is recorded as publicly disclosed in company proxy statements. Financial data deployed within the Diligent Intel platform is normalized and based on information provided by Capital IQ. The performance metrics used in the analysis are selected by Glass Lewis and standardized across companies by industry. Pertaining to U.S. companies, these metrics may differ from the key metrics disclosed by individual companies to meet SEC pay-versus-performance rules. Diligent Intel is a specialist provider of governance research and data analytics. It provides real time data and powerful analytical tools, for independent analysis of corporate governance practices of leading listed companies across the globe, in a single convenient solution. Diligent Corporation and/or its affiliates and suppliers do not make any representation or warranty, express or implied, of any nature, and do not accept any responsibility or liability of any kind, including with respect to the accuracy, completeness or suitability for any purpose of the information contained herein arising from the use of the Diligent Intel platform in connection with this Proxy Paper in any manner whatsoever.

COMPANY UPDATES

EXEMPT SOLICITATION

On November 8, 2024, the National Legal and Policy Center ("NLPC") filed an [exempt solicitation](#) urging shareholders to vote against Reid Hoffman as a nominee to the board. The NLPC cited the following items as immediate concerns regarding this director:

- Dubious associations and questionable judgment; and
- Abrasive, off-putting campaign activism offends across the political spectrum.

CYBERSECURITY INCIDENTS

As discussed in last year's Proxy Paper, on July 11, 2023, the Company published a series of blog posts (found [here](#) and [here](#)) detailing a cyberattack that allowed a threat actor tracked as Storm-0558 to gain unauthorized access to certain customer emails. The Company stated that they detected and mitigated the attack.

On August 11, 2023, the Department of Homeland Security's Cyber Safety Review Board (the "CSRB") [announced](#) that it would conduct a review of the malicious targeting of cloud computing environments, including the July 2023 incident involving the Company.

On April 2, 2024, the CSRB [announced](#) the release of its [findings and recommendations](#) following its independent review of the July 2023 incident. In the review, the CSRB detailed the Company's operational and strategic decisions that led to the intrusion and provided recommendations for specific practices for industry and government to implement to prevent a future intrusion of this magnitude.

In its report, the CSRB recommended that the Company develop and publicly share a plan with specific timelines to make fundamental, security-focused reforms across the Company and its suite of products. According to the CSRB, the Company fully cooperated with the CSRB's review.

On May 3, 2024, the Company published a [blog post](#) detailing the actions the Company had taken and plans to take in response to the report's findings. On September 23, 2024, the Company published a subsequent [blog post](#) to provide a progress update on its security initiatives.

In addition, in January and March 2024, the Company disclosed that it detected a separate cyberattack referred to as the "Midnight Blizzard" attack in the CSRB's report and on the Company's blog. The Company disclosed that beginning in late November 2023, a nation-state-associated threat actor gained access to and exfiltrated information from certain employee email accounts, including members of the Company's senior leadership team, and employees in its cybersecurity and legal teams. The Company disclosed that it removed the threat actor's access to the email accounts in January 2024.

We will continue to monitor this issue going forward.

1.00: ELECTION OF DIRECTORS

SPLIT

PROPOSAL REQUEST: Election of twelve directors

ELECTION METHOD: Majority

RECOMMENDATIONS & CONCERNS:

AGAINST: H. Johnston (Serves on too many boards)

FOR: S. Nadella ; R. Hoffman ; T. List ; C. MacGregor ; M. Mason ; S. Peterson ; P. Pritzker ; C. Rodriguez ; C. Scharf ; J. Stanton ; E. Walmsley

PROPOSAL SUMMARY

Shareholders are being asked to elect 12 nominees to each serve a one-year term.

BOARD OF DIRECTORS

UP	NAME	AGE	GENDER	DIVERSE ⁺	GLASS LEWIS CLASSIFICATION	COMPANY CLASSIFICATION	OWN ^{**}	COMMITTEES					TERM START	TERM END	YEARS ON BOARD
								AUDIT	COMP	GOV	NOM	E&S [^]			
✓	Satya Nadella* ·CEO ·Chair	57	M	Yes	Insider ¹	Not Independent	Yes						2014	2024	10
✓	Reid G. Hoffman	57	M	No	Independent ²	Independent	Yes					✓	2017	2024	7
✗	✓ Hugh F. Johnston*	63	M	No	Independent	Independent	Yes	C ^x					2017	2024	7
✓	Teri L. List	61	F	No	Independent	Independent	Yes	✓ ^x		✓	✓		2014	2024	10
✓	Catherine MacGregor*	52	F	Yes	Independent	Independent	Yes					✓	2023	2024	1
✓	Mark Mason*	55	M	No	Independent	Independent	Yes			✓	✓		2023	2024	1
✓	Sandra E. Peterson ·Lead Director	65	F	No	Independent ³	Independent	Yes		✓	C	C		2015	2024	9
✓	Penny S. Pritzker	65	F	No	Independent	Independent	Yes					C	2017	2024	7
✓	Carlos A. Rodriguez	60	M	Yes	Independent	Independent	Yes	✓ ^x	C				2021	2024	3
✓	Charles W. Scharf*	59	M	No	Independent	Independent	Yes		✓	✓	✓		2014	2024	10
✓	John W. Stanton	69	M	No	Independent	Independent	Yes	✓				✓	2014	2024	10
✓	Emma N. Walmsley*	55	F	No	Independent	Independent	Yes		✓			✓	2019	2024	5

C = Chair, * = Public Company Executive, ^x = Audit Financial Expert, ✗ = Withhold or Against Recommendation

1. Chair and CEO.
2. Co-founder of LinkedIn, which was acquired by the Company in December 2016.
3. Lead independent director.

⁺Reflects racial/ethnic diversity reported either by the Company or by another company where the individual serves as a director. Only racial/ethnic diversity reported by the Company will be reflected in the Company's reported racial/ethnic board diversity percentage listed elsewhere in this Proxy Paper, if available.

^{**}Percentages displayed for ownership above 5%, when available

[^]Indicates board oversight responsibility for environmental and social issues. If this column is empty, it indicates that this responsibility hasn't been formally designated and codified in committee charters or other governing documents.

NAME	ATTENDED AT LEAST 75% OF MEETINGS	PUBLIC COMPANY EXECUTIVE	ADDITIONAL PUBLIC COMPANY DIRECTORSHIPS
Satya Nadella	Yes	Yes	None
Reid G. Hoffman	Yes	No	(2) Joby Aviation, Inc. ; Aurora Innovation, Inc.
Hugh F. Johnston	Yes	Yes	(1) HCA Healthcare, Inc.
Teri L. List	Yes	No	(3) Danaher Corporation ; Visa Inc. ; Lululemon Athletica Inc.
Catherine MacGregor	Yes	Yes	(1) Engie ^E
Mark Mason	Yes	Yes	None
Sandra E. Peterson	Yes	No	None
Penny S. Pritzker	Yes	No	None
Carlos A. Rodriguez	Yes	No	(1) Automatic Data Processing, Inc.
Charles W. Scharf	Yes	Yes	(1) Wells Fargo & Company ^E
John W. Stanton	Yes	No	(1) Costco Wholesale Corporation
Emma N. Walmsley	Yes	Yes	(1) GSK plc ^E

E = Executive

MARKET PRACTICE

BOARD	REQUIREMENT	BEST PRACTICE	2022*	2023*	2024*
Independent Chair	No ¹	Yes ⁶	No	No	No
Board Independence	Majority ²	66.7% ⁶	92%	92%	92%
Gender Diversity	N/A ⁵	N/A ⁵	41.7%	41.7%	41.7%
COMMITTEES	REQUIREMENT	BEST PRACTICE	2022*	2023*	2024*
Audit Committee Independence	100% ³	100% ⁶	100%	100%	100%
Independent Audit Chair	Yes ³	Yes ⁶	Yes	Yes	Yes
Compensation Committee Independence	100% ⁴	100% ⁶	100%	100%	100%
Independent Compensation Chair	Yes ⁴	Yes ⁶	Yes	Yes	Yes
Nominating Committee Independence	100% ⁴	100% ⁶	100%	100%	100%
Independent Nominating Chair	Yes ⁴	Yes ⁶	Yes	Yes	Yes

* Based on Glass Lewis classification

1. Nasdaq Corporate Governance Requirements
2. Independence as defined by Nasdaq listing rules
3. Securities Exchange Act Rule 10A-3 and Nasdaq listing rules

4. Non-independent member allowed under certain circumstances in Nasdaq listing rules
5. No current marketplace listing requirement
6. CII

Glass Lewis believes that boards should: (i) be at least two-thirds independent; (ii) have standing compensation and nomination committees comprised solely of independent directors; and (iii) designate an independent chair, or failing that, a lead independent director.

GLASS LEWIS ANALYSIS

We believe it is important for shareholders to be mindful of the following:

POTENTIAL OVERCOMMITMENT

Shareholders should be mindful of the following commitment levels:

DIRECTOR	ROLE AT THE COMPANY	OUTSIDE PUBLIC COMPANY DIRECTORSHIPS	ROLE
Hugh F. Johnston	NED	HCA Healthcare, Inc.	NED
		N/A	Executive at a public company where they do not serve as a director

DIVERSITY POLICIES AND DISCLOSURE

FEATURE	COMPANY DISCLOSURE
Director Race and Ethnicity Disclosure	Aggregate
Diversity Considerations for Director Candidates	Gender and race/ethnicity
"Rooney Rule" or Equivalent	Yes
Director Skills Disclosure (Tabular)	Matrix
*Overall Rating: Exemplary	
Company-Reported Percentage of Racial/Ethnic Minorities on Board: 25.0%	

*For more information, including detailed explanations of how Glass Lewis assesses these features, please see Glass Lewis' [Approach to Diversity Disclosure Ratings](#).

The Company has provided exemplary disclosure of its board diversity policies and considerations. Areas to potentially improve this disclosure are as follows:

Race and Ethnicity Disclosure - The Company has not disclosed the racial/ethnic diversity of directors in a way that is both delineated from other diversity measures and on an individual basis. Glass Lewis believes that shareholders benefit from clear disclosure of racial/ethnic board diversity on an individual basis.

DIRECTOR COMMITMENTS

Director List serves as a member of the audit committee while serving on three additional public company audit committees. Ordinarily, we believe that the time commitment required by service on this many audit committees may preclude a director from dedicating the time and attention required to fulfill their duty to this Company's shareholders. However, we note that this director is a former CFO of The Gap, Inc., Dick's Sporting Goods, Inc., and Kraft Foods Group, Inc. In light of this level of experience and professional training, we believe this director's simultaneous service on four public company audit committees is acceptable.

BOARD SKILLS

Glass Lewis believes that depth and breadth of experience is crucial to a properly functioning board. We believe shareholders' interests are best served when boards proactively address a lack of diversity through targeted refreshment, linking organic succession planning with the skill sets required to guide and challenge management's implementation of the board's strategy.

We have reviewed the non-employee directors' current mix of skills and experience as follows*:

BASIC INFORMATION				CORE SKILLS							SECTOR-SPECIFIC SKILLS				
DIRECTOR	AGE	GENDER	TENURE	CORE	FINANCE/ RISK	LEGAL/ POLICY	SENIOR EXEC	CYBER/ IT	E&S	HCM	INTL SALES/ MRKTS	TECH/ ENG	MFG/ SCM/ GLOBAL OPS	COMMS/ MRKTING/ E-COM	TECH
Reid G. Hoffman	57	M	7	X			X			X					X
Hugh F. Johnston	63	M	7		X		X				X		X		
Teri L. List	61	F	10		X		X						X		
Catherine MacGregor	52	F	1				X		X	X	X	X	X		
Mark A. L. Mason	55	M	1		X		X				X				
Sandra E. Peterson	65	F	9				X			X	X		X	X	
Penny S. Pritzker	65	F	7			X			X						
Carlos A. Rodriguez	60	M	3	X	X		X			X					X
Charles W. Scharf	59	M	10		X		X			X					
John W. Stanton	69	M	10				X			X				X	X
Emma N. Walmsley	55	F	5				X			X	X		X	X	

*Please note that the above information is for guidance only and has been compiled using the Company's most recent disclosure and/or additional public sources as necessary. It is not intended to be exhaustive. For further information, please refer to the Glass Lewis [Board Skills Appendix](#).

RECOMMENDATIONS

We recommend that shareholders oppose the election of the nominee listed below based on the following:

DIRECTOR COMMITMENTS

Nominee **JOHNSTON** serves as senior executive vice president and CFO of The Walt Disney Company, a public company, while serving on a total of two public company boards. We believe that the time commitment required by this number of external board memberships, in conjunction with executive duties, may preclude this nominee from dedicating the time necessary to fulfill the responsibilities required of directors.

We do not believe there are substantial issues for shareholder concern as to any other nominee.

We recommend that shareholders vote:

AGAINST: Johnston

FOR: Hoffman; List; MacGregor; Mason; Nadella; Peterson; Pritzker; Rodriguez; Scharf; Stanton; Walmsley

2.00: ADVISORY VOTE ON EXECUTIVE COMPENSATION

FOR

PROPOSAL REQUEST:	Approval of Executive Pay Package	PAY FOR PERFORMANCE GRADES:	FY 2024 C FY 2023 C FY 2022 C
PRIOR YEAR VOTE RESULT (FOR):	93.3%	RECOMMENDATION:	FOR
STRUCTURE:	Fair		
DISCLOSURE:	Fair		

EXECUTIVE SUMMARY

SUMMARY ANALYSIS

The pay program, while not without fault, continues to strongly align pay with performance. At present, we believe shareholder support of the proposal is reasonable.

COMPENSATION HIGHLIGHTS

- STI: Performance-based;
 - Below target payouts for the CEO following the committee's application of negative discretion (his payout was reduced by over 50%); and
 - Above target payouts for non-CEO NEOs
- LTI: Performance-based and time-based; FY2022-FY2024 performance cycle paid out above target in aggregate
- One-time: None granted during the past fiscal year

SUMMARY COMPENSATION TABLE

NAMED EXECUTIVE OFFICERS	BASE SALARY	BONUS & NEIP	EQUITY AWARDS	TOTAL COMP
Satya Nadella <i>Chairman and Chief Executive Officer</i>	\$2,500,000	\$5,200,000	\$71,236,392	\$79,106,183
Amy E. Hood <i>Executive Vice President and Chief Financial Officer</i>	\$1,000,000	\$3,642,750	\$21,094,956	\$25,799,206
Judson B. Althoff <i>Executive Vice President and Chief Commercial Officer</i>	\$960,000	\$3,497,040	\$18,515,353	\$23,050,327
Bradford L. Smith <i>Vice Chair and President</i>	\$1,000,000	\$3,642,750	\$18,684,175	\$23,439,793
Christopher D. Young <i>Executive Vice President, Business Development, Strategy, and Ventures</i>	\$850,000	\$2,023,680	\$9,040,931	\$12,034,703
CEO to Avg NEO Pay:				3.75: 1

CEO SUMMARY

	2024 SATYA NADELLA	2023 SATYA NADELLA	2022 SATYA NADELLA
Total CEO Compensation	\$79,106,183	\$48,512,537	\$54,946,310
1-year TSR	32.3%	33.9%	-4.4%
CEO to Peer Median *	2.6:1	1.8:1	1.9:1
Fixed/Perf.-Based/Discretionary **	4.6% / 95.4% / 0.0%	4.8% / 95.2% / 0.0%	4.2% / 95.8% / 0.0%

* Calculated using the first Company-disclosed peer group. ** Percentages based on the CEO Compensation Breakdown values.

CEO COMPENSATION BREAKDOWN

FIXED	Cash		\$2.7M
	Salary		\$2.5M
	Benefits / Other		\$169,791
	Total Fixed		\$2.7M
PERFORMANCE-BASED	Performance Shares		\$50.0M*
	Long-term Incentive Plan		\$50.0M
	Target/Maximum	152,551 shares / 457,653 shares	
	Metrics	LinkedIn Sessions, Search and News Advertising Revenue (ex TAC) Growth, TSR, Azure and Other Cloud Services Revenue Growth, Xbox Content and Services Revenue Growth, Microsoft Cloud Revenue Growth (Excluding Azure and Other Cloud Services)	
	Performance Periods	See award terms below	
	Additional Vesting / Deferral Period	See award terms below	
	Cash		\$5.2M
	Short-term Incentive Plan		\$5.2M
	Target/Maximum	\$7.5M / \$15.0M	
	Metrics	Customers and Stakeholders, Non-GAAP Revenue, Non-GAAP Operating Income, Product and Strategy, Culture, Diversity and Sustainability	
	Performance Period	1 year	
	Additional Vesting / Deferral Period	-	
Total Performance-Based			\$55.2M
Awarded Incentive Pay			\$55.2M
Total Pay Excluding change in pension value and NQDCE			\$57.9M

*Reflects the aggregate value of performance shares granted for the FY2024-FY2026 performance cycle.

PEER GROUP REVIEW ^{1 2 3 4}

THE COMPANY USES TWO PEER GROUPS FOR SETTING PAY LEVELS.

TECHNOLOGY PEER GROUP

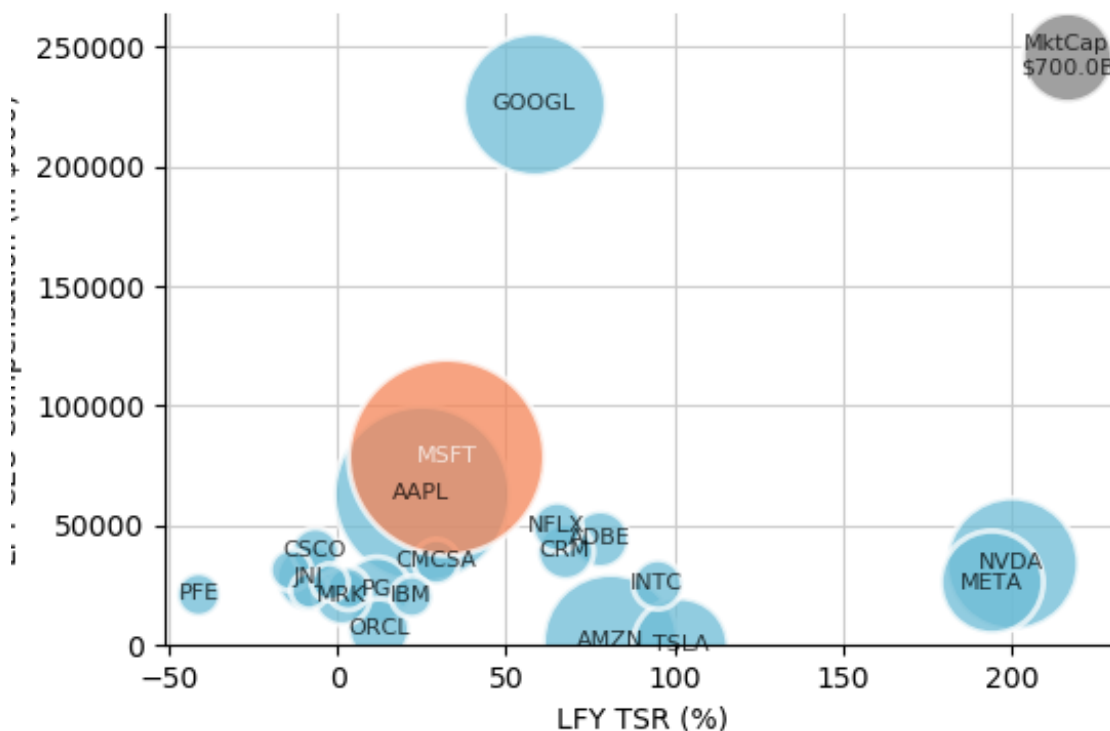
This peer group consists of 12 companies. Total NEO compensation is not benchmarked to a specific percentile of this peer group.

	MARKET CAP	REVENUE	CEO COMP	1-YEAR TSR	3-YEAR TSR	5-YEAR TSR
75th PERCENTILE OF PEER GROUP	\$1538.8B	\$221.1B	\$42.3M	87.7%	16.1%	21.9%
MEDIAN OF PEER GROUP	\$300.5B	\$57.6B	\$30.5M	62.8%	8.8%	17.3%
25th PERCENTILE OF PEER GROUP	\$202.4B	\$44.4B	\$21.9M	17.0%	2.4%	11.9%
COMPANY	\$3321.9B (Highest)	\$245.1B (76th %ile)	\$79.1M (88th %ile)	32.3% (39th %ile)	19.2% (87th %ile)	28.5% (87th %ile)

GENERAL INDUSTRY PEER GROUP

This peer group consists of 11 companies. Total NEO compensation is not benchmarked to a specific percentile of this peer group.

	MARKET CAP	REVENUE	CEO COMP	1-YEAR TSR	3-YEAR TSR	5-YEAR TSR
75th PERCENTILE OF PEER GROUP	\$377.3B	\$121.6B	\$31.6M	29.1%	2.6%	11.8%
MEDIAN OF PEER GROUP	\$194.8B	\$87.0B	\$25.3M	1.9%	-2.6%	7.2%
25th PERCENTILE OF PEER GROUP	\$158.5B	\$60.1B	\$21.6M	-8.6%	-4.3%	-2.9%
COMPANY	\$3321.9B (Highest)	\$245.1B (Highest)	\$79.1M (Highest)	32.3% (75th %ile)	19.2% (Highest)	28.5% (88th %ile)



¹ Market capitalization figures are as of fiscal year end dates. Source: Capital IQ

² Annual revenue figures are as of fiscal year end dates. Source: Capital IQ

³ Annualized TSR figures are as of fiscal year end dates. Source: Capital IQ

⁴ Annual CEO compensation data based on the most recent proxy statement for each company.

EXECUTIVE COMPENSATION STRUCTURE - SYNOPSIS

FIXED

Base salaries did not increase significantly during the past fiscal year.

SHORT-TERM INCENTIVES

STI PLAN

AWARDS GRANTED (PAST FY)	<i>Cash</i>
TARGET PAYOUTS	<i>\$7,500,000 for the CEO and up to \$2,500,000 for the other NEOs</i>
MAXIMUM PAYOUTS	<i>\$15,000,000 for the CEO and up to \$5,000,000 for the other NEOs</i>
ACTUAL PAYOUTS	<i>\$5,200,000 for the CEO and up to \$3,642,750 for the other NEOs</i>

Performance is measured over one year.

The metric weightings below apply to the CEO's payout formula. For all other NEOs, payouts are weighted 50% financial metrics and 50% operational metrics.

The CEO's payout as determined by the plan was \$10.66 million, or 142% of target. However, the Company states that given the focus and speed required for the changes that today's cybersecurity threat landscape showed were necessary, it would be appropriate to reduce the CEO's payout by more than 50%. As such, his final payout was \$5.2 million or approximately 69% of target. Payouts for non-CEO NEOs remained at the level determined by the plan, which was above target (119% to 146% of target).

METRICS FOR OPERATIONAL PERFORMANCE (30%)	PRODUCT AND STRATEGY	CUSTOMERS AND STAKEHOLDERS	CULTURE, DIVERSITY AND SUSTAINABILITY
	Absolute	Absolute	Absolute
	Weighting	10%	10%
	Actual Performance	150% of target	160% of target
		145% of target	

METRICS FOR FINANCIAL PERFORMANCE (70%)	NON-GAAP OPERATING INCOME		NON-GAAP REVENUE	
		Absolute		Absolute
	Weighting	35%		35%
	Threshold Performance	\$94.5B		\$228.9B
	Target Performance	\$103.3B		\$243.5B
	Maximum Performance	\$116.7B		\$256.9B
	Actual Performance	\$110.9B		\$246.1B

LONG-TERM INCENTIVES

LTI PLAN

AWARDS GRANTED (PAST FY)	<i>Performance shares and restricted stock</i>
TARGET PAYOUTS*	<i>Performance Shares: 152,551 shares for the CEO and up to 26,697 shares for the other NEOs</i>
MAXIMUM PAYOUTS*	<i>Performance Shares: 457,653 shares for the CEO and up to 80,091 shares for the other NEOs</i>
TIME-VESTING PAYOUTS	<i>Restricted Stock: Up to 26,697 shares for the non-CEO NEOs</i>

Performance is measured over three one-year periods for weighted metrics (goals set annually) and over a full three-year period for relative TSR. Earned awards vest after the full three-year period.

Restricted stock awards vest over four years.

The TSR modifier is measured relative to the S&P 500. The modifier will not increase the shares earned for a performance period unless absolute TSR for the performance period is positive.

*The target and maximum payouts noted above reflect payout potential for the entire three-year performance period (FY2024 to FY2026).

METRICS FOR PERFORMANCE SHARES MODIFIER	TSR	
	Relative	
	Weighting	Modifier (0.75x to 1.5x)
	Threshold Performance	20th %ile
	Target Performance	40th %ile to 60th %ile
	Maximum Performance	80th %ile

METRICS FOR PERFORMANCE SHARES (2024 TRANCHE)	AZURE & OTHER CLOUD SERVICES REVENUE GROWTH	MICROSOFT CLOUD REVENUE GROWTH (EXCLUDING AZURE & OTHER CLOUD SERVICES)	LINKEDIN SESSIONS	SEARCH & NEWS ADVERTISING REVENUE (EX TAC) GROWTH	XBOX CONTENT & SERVICES REVENUE GROWTH
	Absolute	Absolute	Absolute	Absolute	Absolute
	Weighting	35%	35%	10%	10%
	Threshold Performance	18.92%	7.65%	76.78	-0.26%
	Target Performance	25.89%	15.4%	82.56	10.82%
	Maximum Performance	32.88%	22.82%	88.34	21.91%
	Actual Performance	29.15%	15.75%	83.97	11.73%

Note: Allocation is at least 60% based on Glass Lewis calculations.

RISK-MITIGATING POLICIES

CLAWBACK POLICY	Yes - expanded policy (not strictly restatement-dependent)
ANTI-HEDGING POLICY	Yes
STOCK OWNERSHIP GUIDELINES	Yes - all NEOs

SEPARATION & CIC BENEFITS

HIGHEST SEVERANCE ENTITLEMENT	1x base salary and bonus
CIC EQUITY TREATMENT	Double-trigger acceleration
EXCISE TAX GROSS-UPS	No

OTHER FEATURES

LFY CEO TO MEDIAN EMPLOYEE PAY RATIO	408:1*
E&S METRICS FOR THE CEO	Environment, Human Capital Management, Diversity and Climate
BENCHMARK FOR CEO PAY	No specific benchmark

*The Company-disclosed median employee pay for the year in review was \$193,744.

OTHER COMPENSATION DISCLOSURES

COMPENSATION ACTUALLY PAID (YEAR-END CEO)	\$171,304,590 for FY2024 and \$93,155,708 for the prior fiscal year
REPORTED TSR*	\$227 for FY2024 and \$172 for the prior fiscal year
KEY PVP METRICS	Non-GAAP revenue, relative TSR as compared to the S&P 500, Azure and Other Cloud Services revenue, Microsoft Cloud revenue (excluding Azure and Other Cloud Services) and non-GAAP operating income

*Reported TSR reflects the year-end value of an initial fixed \$100 investment at the start of the required reporting period under SEC Pay Vs Performance (PVP) disclosure rules.

GLASS LEWIS ANALYSIS

This proposal seeks shareholder approval of a non-binding, advisory vote on the Company's executive compensation. Glass Lewis believes firms should fully disclose and explain all aspects of their executives' compensation in such a way that shareholders can comprehend and analyze the company's policies and procedures. In completing our assessment, we consider, among other factors, the appropriateness of performance targets and metrics, how such goals and metrics are used to improve Company performance, the peer group against which the Company believes it is competing, whether incentive schemes encourage prudent risk management and the board's adherence to market best practices. Furthermore, we also emphasize and evaluate the extent to which the Company links executive pay with performance.

PROGRAM FEATURES ¹

POSITIVE

- Alignment of pay with performance
- LTIP performance-based
- STIP performance-based
- STI-LTI payout balance
- No single-trigger CIC benefits
- Anti-hedging policy
- Enhanced clawback policy for NEOs
- Executive stock ownership guidelines for NEOs

NEGATIVE

- Short performance period under LTIP

¹ Both positive and negative compensation features are ranked according to Glass Lewis' view of their importance or severity

AREAS OF FOCUS

VARIABLE COMPENSATION

Vesting Below Median

Policy Perspective: Long-term incentive plans that allow for significant payouts for below-median performance effectively may reward NEOs for significant underperformance. Shareholders may question whether such structures are fully appropriate.

Analyst Comment: We maintain concerns about the limited effect of downside modification resulting from TSR underperformance. By way of example, TSR performance as low as the 20th percentile would only reduce payouts by 25%, which could result in high payouts despite performance significantly below the median. On the other hand, the upside modification offers an opportunity of up to 50% for performance at or above the 80th percentile. No modification is made to derived payouts as a result of relative TSR performance between the 40th and 60th percentiles.

Performance Period of Long-Term Awards

Policy Perspective: Performance measurement periods of less than three years may fail to sufficiently incentivize long-term thinking among executives and may lead to payouts that are not reflective of a Company's overall health.

High Maximum Payout Limit under STIP

Analyst Comment: Shareholders should be mindful of the high payout limits under the short-term incentive plan. In our view, excessive upside opportunities under a bonus plan may unduly incentivize short-term performance and undermine a long-term focus on company performance among executives. As noted in the prior year, this concern is underscored by the potentially limited emphasis on long-term focus given the highlighted issue with a short-performance period under the LTIP. Here, the CEO's short-term incentive opportunity at target is three times base salary, while the maximum payout under the plan is six times his base salary, which we consider excessive. Be that as it may, in consideration of the application of negative discretion, and the sustained alignment of pay and performance, we do not believe shareholders should be overly worried with regard to this concern.

Total Target Direct Compensation for the CEO

Analyst Comment: Indeed, Mr. Nadella's total target direct compensation hovers near \$60 million, which sits well above the median of the Company's self-disclosed and Glass Lewis-identified peers. However, certain concessions must be made when considering the Company's relative size in terms of market capitalization and revenue, which also place the Company well above the median of these two peer groups. There were also no increases to Mr. Nadella's total target direct compensation during the year in review. At present, our pay for performance model indicates that the Company is providing for executive pay levels commensurate with Company performance, and we do not believe the quantum of pay or the structure of the pay program warrants a negative recommendation.

2024 PAY FOR PERFORMANCE: C

Policy Perspective: "C" grades in the Glass Lewis pay-for-performance model indicate an adequate alignment of pay with performance, where the gap between compensation and performance rankings is not significant.

CONCLUSION

We recommend that shareholders vote **FOR** this proposal.

3.00: RATIFICATION OF AUDITOR

FOR

PROPOSAL REQUEST: Ratification of Deloitte & Touche
PRIOR YEAR VOTE RESULT (FOR): 95.1%
BINDING/ADVISORY: Advisory
REQUIRED TO APPROVE: Majority of votes cast
AUDITOR OPINION: Unqualified

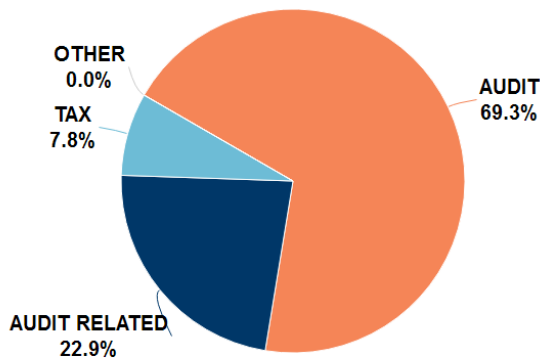
RECOMMENDATIONS & CONCERNS:
FOR- No material concerns

AUDITOR FEES

	2024	2023	2022
Audit Fees:	\$56,280,000	\$45,630,000	\$43,345,000
Audit-Related Fees:	\$18,643,000	\$10,713,000	\$9,617,000
Tax Fees:	\$6,334,000	\$5,049,000	\$4,545,000
All Other Fees:	\$10,000	\$10,000	\$10,000
Total Fees:	\$81,267,000	\$61,402,000	\$57,517,000
Auditor:	Deloitte & Touche	Deloitte & Touche	Deloitte & Touche
1-Year Total Fees Change:		32.4%	
2-Year Total Fees Change:		41.3%	
2024 Fees as % of Revenue*:		0.033%	

* Annual revenue as of most recently reported fiscal year end date. Source: Capital IQ

Years Serving Company:	41
Restatement in Past 12 Months:	No
Alternative Dispute Resolution:	No
Auditor Liability Caps:	No
Lead Audit Partner:	Michael Scott Thompson
Critical Audit Matters:	3
	• Revenue Recognition • Income Taxes – Uncertain Tax Positions • Business Combinations – Estimate for Valuation of Acquired Intangible Assets



GLASS LEWIS ANALYSIS

The fees paid for non-audit-related services are reasonable and the Company discloses appropriate information about these services in its filings.

We recommend that shareholders vote **FOR** the ratification of the appointment of Deloitte & Touche as the Company's auditor for fiscal year 2025.

4.00: SHAREHOLDER PROPOSAL REGARDING RISKS OF DEVELOPING MILITARY WEAPONS

FOR

PROPOSAL REQUEST:	That the board report on the risks of being identified as a company developing weapons used by the military	SHAREHOLDER PROPONENT:	Harrington Investments, Inc.
BINDING/ADVISORY:	Precatory		
PRIOR YEAR VOTE RESULT (FOR):	15.2%	REQUIRED TO APPROVE:	Majority of votes cast
RECOMMENDATIONS, CONCERNS & SUMMARY OF REASONING:			
FOR -	Additional disclosure could help shareholders understand financial and and reputational risks from work with the military		

SASB MATERIALITY

PRIMARY SASB INDUSTRY: Software & IT Services

FINANCIALLY MATERIAL TOPICS:

- Environmental Footprint of Hardware Infrastructure
- Recruiting & Managing a Global, Diverse & Skilled Workforce
- Managing Systemic Risks from Technology Disruptions
- Data Privacy & Freedom of Expression
- Data Security
- Intellectual Property Protection & Competitive Behavior

GLASS LEWIS REASONING

- Given the potential financial, reputational, and human-capital-related risks that could be realized via the Company's work with the military, we believe that a more robust discussion on this issue should be disclosed.

PROPOSAL SUMMARY

Text of Resolution: *BE IT RESOLVED Shareholders request that the board issue an independent, third-party report, at reasonable expense and excluding proprietary information, to assess the reputational and financial risks to the company for being identified as a company involved in the development of weapons used by the military.*

Proponent's Perspective

- The Company developed an augmented reality headset to provide night vision, thermal sensing, and monitoring of vital signs, initially intended for gaming purposes, but then the U.S. Army adapted this product to be used for military training and combat;
- In March 2021, the Company was awarded a \$479 million Integrated Visual Augmentation System ("IVAS") contract with the U.S. Department of the Army, which later became a \$22 billion contract for a semi-custom version of IVAS to rapidly develop, test, and manufacture a single platform that soldiers can use to fight, rehearse, and train that provides increased lethality, mobility, and situational awareness necessary to achieve overmatch against its current and future adversaries;
- In 2019, amid the contract negotiation with the military, Company employees pushed back in a letter to the Company stating they "do not want to become war profiteers" and they "did not sign up to develop weapons" and that they demand a say in how their work is used;
- Regardless of how the Company positions itself by making public statements, the outcome is a significant contract with the U.S. military and the Company is actively working to develop what employees, other stakeholders, including shareholders, and the public view as a weapons system used in war;
- In December, images of Chinese military personnel wearing HoloLens 2 headsets were seen on China's State Media while, simultaneously, U.S. lawmakers urged increased restrictions on exports to stymie the Chinese military's access to U.S. technologies, especially those with dual application in commercial and defense sectors; and
- Involvement in the development of weapons poses a serious risk to a company's reputation, especially for shareholders and stakeholders.

Board's Perspective

- The Company has worked with the U.S. Department of Defense on a longstanding and reliable basis for four decades;
- The Company's senior leadership team made a decision not to withhold technology from institutions that citizens have elected in democracies to protect the freedoms they enjoy;
- The Company is committed to providing its technology to the U.S. military and to providing the Company's expertise and perspective on technology issues ranging from cybersecurity to the ethical use of artificial intelligence ("AI");
- The Company's Responsible AI team and external ethics experts were involved in reviewing the Integrated Visual Augmentation System to help guide its development, and the Company is committed to proactively working to address the ethical issues that new technology creates for the military by engaging in discussions with the country's institutions, including engagement before Congress, engagement with the Executive Branch, and engagement with the military itself; and
- The Company does not ask or expect everyone who works at the Company to support every position it takes, and it respects that some employees may not wish to work on certain projects.

THE PROPONENT

Harrington Investments

[Harrington Investments](#), founded by John Harrington, is an investment manager that is "dedicated to managing portfolios for individuals, foundations, non-profits, and family trusts to maximize financial, social and environmental performance." It states that its strategy involves "inclusionary and exclusionary investment criteria to develop socially responsible portfolios, shareholder advocacy to influence the policies of corporations owned by Harrington Investments and [its] clients, and community investing to provide needed capital for small businesses, job training and vital community services." As of June 30, 2023, Harrington Investments [managed](#) approximately \$316.2 million on a discretionary basis and \$569,711 on a non-discretionary basis.

Harrington [states](#) that it has been "actively engaged in shareholder advocacy for over 35 years" and that, in that time, it has "worked to persuade more than 120 publicly traded corporations to increase their economic, social and environmental accountability and responsibility." It states that its shareholder advocacy uses several tactics including: (i) communicating directing with corporate management and elected and appointed government officials; (ii) filing and co-filing shareholder proposals to "voice [its] concerns and to encourage productive dialogue with corporate management on important social, environmental, and corporate governance issues;" (iii) collaborating and coordinating shareholder advocacy with other investors and national shareholder organizations; (iv) encouraging companies to "initiate innovative policies to further environmental and social responsibility;" and (v) encouraging companies to incorporate codes of conduct for workers in U.S. companies to ensure better human and labor rights.

Based on the disclosure provided by companies concerning the identity of proponents, during the first half of 2023, Harrington Investments submitted ten shareholder proposals that received an average of 15.6% support (excluding abstentions and broker non-votes), with none of its proposals receiving majority support.

GLASS LEWIS ANALYSIS

Glass Lewis recommends that shareholders take a close look at proposals such as this to determine whether the actions requested of the Company will clearly lead to the enhancement or protection of shareholder value. Glass Lewis believes that directors who are conscientiously exercising their fiduciary duties will typically have more and better information about the Company and its situation than shareholders. Those directors are also charged with making business decisions and overseeing management. Our default view, therefore, is that the board and management, absent a suspicion of illegal or unethical conduct, will make decisions that are in the best interests of shareholders.

We believe that it is prudent for firms to actively evaluate risks to shareholder value stemming from global activities and human rights practices. This evaluation is particularly important for those companies with operations in sensitive regions prone to allegations of human rights abuses. As has been seen with other companies, allegations of human rights abuses can inflict, at a minimum, reputational damage on targeted companies and may dramatically affect shareholder value.

In this case, the proposal is centered on the Company's [HoloLens 2](#) product, which is a holographic computer that the Company is using in the development of combat goggles for the military, among other things. Given recent controversies concerning the use of military technologies, the Company's product offering could present significant reputational and human rights-related risks.

CONCERNS REGARDING COMPANY SUPPORT FOR GOVERNMENT AGENCIES

The Company has experienced significant employee backlash in recent years as a result of its work with government agencies. For example, in 2018, employees expressed dissent regarding the Company's decision to contract with U.S. Immigration and Customs Enforcement ("ICE"). Initially, the Company announced that it was "proud to support" the agency, but later condemned its use of family separation. Employees expressed their dissent internally, questioning whether they should stay with the Company or consider leaving (Brian Menegus. "[Microsoft Employees Up in Arms Over Cloud Contract with ICE](#)." *Gizmodo*. June 18, 2018). More recently, reporting has shown that tech companies, including the Company, Amazon, and Alphabet are using tactics to dodge public scrutiny and continue working with U.S. immigration agencies, despite employee backlash and company policies. The companies have secured dozens of cloud contracts with both ICE and CBP, according to *Business Insider*, by using third parties or acting as subcontractors to sell their technology. These contracts were discovered through the use of company names in descriptions, although it is unclear how many other contracts exist that are tied to these companies but do not mention their involvement (Caroline Haskins. "[Google, Amazon, and Microsoft Have Taken Advantage of a Commonly Used but Little-Known Tool to Quietly Enter Dozens of Contracts With ICE and CBP](#)." *Business Insider*. October 4, 2021).

HoloLens Technology

The Company's HoloLens technology has also been a subject of scrutiny from employees. In 2019, the Army showcased

a specially modified HoloLens 2 headsets made by the Company as part of a \$480 million defense contract it had won. The military's version of the headset, known as the Integrated Visual Augmentation System ("IVAS"), is an augmented-reality headset that places digital objects, such as maps or video displays, on top of real world views. In response, approximately 100 employees wrote an [open letter](#) to the Company's CEO demanding that he cancel the IVAS contract and cease the development of any and all weapons technologies. The employees further expressed alarm that the Company was working to provide weapons technology to the U.S. military, helping the government "increase lethality" using tools they built (Todd Haselton. ["How the Army Plans to Use Microsoft's High-Tech HoloLens Goggles on the Battlefield."](#) CNBC. April 6, 2019).

In March 2021, the Company won another contract for over 120,000 headsets, worth up to \$21.88 billion over 10 years. The Company's CEO defended its augmented reality project, stating that the Company would not "withhold technology from institutions that we have elected in democracies to protect the freedoms we enjoy" (Jordan Novet. ["Microsoft Wins U.S. Army Contract for Augmented Reality Headsets. Worth up to \\$21.9 Billion Over 10 Years."](#) CNBC. March 31, 2021).

Early 2023 saw the Company laying off as many as 10,000 workers and reducing its HoloLens business after Congress rejected the Army's request to pay \$400 million for as many as 6,900 of the IVAS goggles, partially in response to problems in the product's testing (Dina Bass. ["Microsoft Job Cuts Hit HoloLens Unit After Setback on Army Goggles."](#) Bloomberg. January 18, 2023). In response to the Army's feedback that its IVAS product was causing nausea and dizziness, however, the Company was able to develop an improved version of the combat goggles, which it was on track to deliver for intensive soldier testing in August 2023. If this improved version proved reliable and provided low-light performance without the episodes of nausea and dizziness soldiers experienced with the previous prototype, the Army reportedly could deploy the devices by 2025; otherwise, the Company faced the possibility that the program could be canceled (Anthony Capaccio. ["Microsoft Poised to Deliver Improved Combat Goggles. US Army Says."](#) Bloomberg. July 19, 2023).

In September of the same year, the Army reported that it had tested 20 prototypes for "weapons compatibility, navigation tools, and mission planning work" and was ready to move on to the next evaluation phase, an 18-month period during which time the teams involved with the HoloLens program would work on making the system "producible and affordable" as well as determining the number of capabilities and levels users would be able to access (Todd South. ["Army Approves Next Phase for Augmented Reality Device."](#) Army Times. September 7, 2023).

Despite the U.S. imposing export controls on advanced microchips in 2022 and tightening of those export controls in 2023, Chinese state media released a short demonstration video in December 2023 showing a member of the Chinese People's Liberation Army using the Company's HoloLens 2 smart glasses while being trained in aircraft maintenance. Meanwhile, U.S. Commerce Secretary Gina Raimondo discussed preparing additional technological curbs on China, with Republican House committee leaders Mike Gallagher and Michael McCaul informing Raimondo that further resources to her department would hinge upon institutional reforms limiting the flow of U.S. technology to American adversaries (Aadil Brar. ["China's State Media Shows Military Using Microsoft's HoloLens 2 Headsets."](#) Newsweek. December 14, 2023).

Additionally, in December 2023, the U.S. Army reported that it performed weapons compatibility checks on the more ruggedized and field-worthy IVAS 1.2 prototypes. The Army explained that 80 more devices were scheduled for delivery in 2024 and 200 more in 2025, with fielding planned for the devices in 2025. Earlier versions of the device, the 1.0 and 1.1 IVAS, would be used for training units and schoolhouses to enable soldiers to tinker with and learn the best applications for the technology (Todd South. ["Army's Mixed Reality Device Nears Fielding with Final Testing in 2024."](#) Army Times. December 29, 2023).

As of October 2024, the Army stated that it was focused on the "extensibility" of the HoloLens device, meaning its ability to control other devices like the soldier-borne sensor microdrone already used by soldiers. Moreover, based on ongoing feedback from testing, the Army was planning several design changes to IVAS 1.2, including: (i) a low-light camera with increased sensitivity; (ii) an improved low-light focus mechanism, compatible with wearing gloves; (iii) more robust bumpers, cables, bungees, and tethered solar caps; (iv) hinge improvements for usability, display clarity, and durability; (v) improved transport case and mission bag for better storage and protections; (vi) minor visor and display improvements for greater clarity and durability; and (vii) software improvements. Soldiers at Fort Campbell in Kentucky were expected to test updates in January 2025, and following that user assessment, soldiers at Fort Carson in Colorado were slated to conduct an operational demonstration with the device in the spring of 2025. Officials stated that once the IVAS 1.2 engineering changes were made and the Army completed the battalion operational assessment and any other testing, the Army could award a production contract. Officials also stated that the development team had fixed the early issues with moisture in the devices, display glitches, and the dizziness reported by users (Todd South. ["Army's Mixed Reality Device Set for Upgrades and Battalion Assessment."](#) Defense News. October 15, 2024).

Meanwhile, the defense technology firm Anduril announced in September 2024 that it had incorporated its Lattice solution into IVAS 1.1 and 1.2 for the Army, enabling it to employ its situational awareness tool that uses AI, computer vision, edge computing, and sensor fusion to detect, track, and classify objects of interest for users. Earlier in the year, the Army

requested \$255 million from Congress for fiscal 2025 towards the purchase of more IVAS systems, as well as \$98 million for research, development, testing, and evaluation related to the technology. Although the Army was planning to transition IVAS to a major capability acquisition pathway by October 2025, it explained that it had not yet determined whether the Company would remain the prime contractor for the program as officials considered the next stage of the modernization effort. Should the program continue to move forward, it could be worth as much as \$21.9 billion (Jon Harper. "[Anduril Integrates AI Tech into Army IVAS Headsets](#)." *DefenseScoop*. September 19, 2024).

COMPANY POLICIES AND DISCLOSURE

In its [Human Rights Annual Report](#), the Company discusses risks and considerations related to its technology. It states that the Company is committed to respecting human rights in its development and deployment of generative artificial intelligence ("AI") technologies. It notes that it has been refining its approach to the development, testing, deployment, and responsible use of AI systems since 2017. Further, as the UN Guiding Principles on Business and Human Rights call for human rights due diligence and exploration of potential adverse human rights impacts and the development of mitigation measures, the Company has applied this approach in its development and deployment of AI technologies. It then affirms that it commissioned a human rights impacted assessment ("HRIA") of AI in 2018 and engaged stakeholders to understand better how to meet the Company's responsibility to respect human rights (p.6). The Company discloses its [six Responsible AI principles](#), which are [based](#) on learnings from its 2018 HRIA. The Company adds that it established its Office of Responsible AI ("ORA") in 2019, to manage its sensitive uses review program (p.6). The Company explains that ORA is a reporting, review, and guidance framework that is triggered when Company personnel are involved in developing or deploying an AI system and the foreseeable use or misuse of that AI system could:

- Have a consequential impact on a user's legal status or life opportunities;
- Present the risk of significant physical or psychological injury; or
- Restrict, infringe upon, or undermine the ability to realize an individual's human rights.

The Company [explains](#) that ORA developed the Responsible AI Standard framework for implementing AI principles across the Company (p.6). It notes that after receiving feedback from diverse disciplines, the Company published an [updated version](#) of the Standard in 2022, [defining](#) requirements for designing, building, and testing AI systems responsibly (p.6). It then provides a number of publicly available resources that encourage responsible AI through the stages of innovation. It also discusses its engagements and its commitment to supporting greater international collaboration on AI governance and multi-stakeholder approaches (pp.7-8).

In its most recent [impact summary](#), the Company lists the following goals related to its technology:

- Continue upholding its AI principles for the responsible development and use of AI within the Company;
- Work with industry leaders and governments to develop new standards for highly capable foundation models;
- Empower the Company's customers and partners to develop and use AI responsibly;
- Invest \$20 billion over five years, starting in 2021, to advance its security solutions, including \$150 million to help U.S. government agencies upgrade protection;
- Provide global threat intelligence, expert guidance, and innovative solutions to help customers, partners, and governments improve their cyber resiliency;
- Preserve customers' control over their data and ability to make informed choices that protect their privacy, while advocating for strong global privacy and data protection laws requiring companies, including the Company, to only collect and use personal data in responsible, accountable ways;
- Protect users from illegal and harmful content and conduct, while respecting human rights; and
- Supporting multistakeholder approaches to address complex, whole-of-society digital safety challenges.

(pp.10-12)

The Company also addresses responsible technology use in its [Global Human Rights Statement](#). Further, it [provides](#) a variety of transparency reports, including a [U.S. National Security Orders report](#), as well as [privacy reports](#).

In October 2021, the Company [released](#) a statement agreeing to conduct additional human rights due diligence regarding the role of its technology and its potential impact on certain communities in select situations. It states that its decision to pursue such investigations stems from guidance from the UN Guiding Principles on Business and Human Rights. Additionally, the Company states that the assessment will help to identify, understand, assess, and address actual or potential adverse human rights impacts of its products and services and business relationships with regard to law enforcement, immigration enforcement, and other government contracts. Moreover, it will include consultation with communities most impacted by the Company's surveillance products, law enforcement, and government contracts (Dina Bass. "[Microsoft Agrees to Human Rights Review of Deals With Law Enforcement and Government](#)." *Time*. October 13, 2021).

The Company also published its 2023 [Human Rights Impact Assessment of Microsoft's Enterprise Cloud and AI Technologies Licensed to U.S. Law Enforcement Agencies](#), which was conducted by members of the Global Business &

Human Rights Practice of the law firm Foley Hoag LLP and which briefly discusses the Company's business relationships with the military. The assessment points out that the 2021 shareholder resolution that prompted the HRIA referenced the Company's contracts with the U.S. military, and it explains that shareholders expressed interest in having the HRIA address sales to all government agencies, both civilian and military. It explains, however, that the Company expressed concern that encompassing the military and civilian dimensions would render the scope of the HRIA too broad to give the subject the "nuanced treatment it deserves." The assessors detail discussing the exclusion of military contracts from the assessment with approximately 20 representatives of civil society organizations involved with the development of the shareholder resolution, and they "voiced disappointment over the decision to exclude military uses, emphasizing that a review of these contracts was expressly called for in the shareholder resolution." Additionally, the assessment explains that in one-on-one interviews, several representatives opined that "the military's use of Microsoft products is intertwined with uses by civilian authorities" and others "noted that Microsoft appears to have significant contracts with the U.S. military, and that the products it provides in military settings stand at substantial risk of furthering serious human rights violations through their potential use by foreign governments to commit atrocities, genocide, and other crimes against humanity" (pp.7-8). In its recommendation regarding human rights due diligence, the HRIA states:

Microsoft should consider conducting theme-specific due diligence exercises or additional HRIAs focused on: assessing the company's effectiveness in implementing the remediative and mitigating steps proposed in this HRIA's recommendations; assessing the company's commercial relationships with military agencies and their impacts on BIPOC and other vulnerable communities; and other human rights challenges prioritized by the stakeholder organizations with which it engages.

(p.50) Looking to the HRIA's key findings, which do not address the Company's contracts with the military, the assessment [states](#) that "in no instance did the Assessors find evidence that Microsoft causes human rights harms." It adds that in most cases, the Company is an "upstream provider of platforms that exist independently of the software and applications that developers create and use on those platforms," which "attenuates any responsibility Microsoft might have for downstream adverse human rights impacts, such as discriminatory policing, surveillance, and incarceration." The assessment clarifies, however, that "the precise degree of Microsoft's responsibility" is "unclear under the UNGPs." Further, the assessment explains that in the few cases where the Company is actively involved in developing products through its consulting services, "the relationship between Microsoft and any downstream adverse impact is more concrete, and Microsoft has at least the responsibility to mitigate impacts." The report's final key finding states that the Company's human rights policies and practices are "robust" but highlights the perception particularly from civil society groups, particularly representatives of the human rights community, that there is "a lack of transparency with respect to product design and deployment." According to the HRIA, "this hinders Microsoft's ability to speak with authority on human rights issues" (pp.1-2).

In May 2023, the Company published [Governing AI: A Blueprint for the Future](#), which examines the Company's legal and regulatory blueprint for the future and discusses its approach to building AI systems that benefit society. However, the document does not directly discuss the Company's technology in relation to military use, but it addresses implementing and building upon new government-led AI safety frameworks (p.5). The Company also realizes that some will seek to use AI as a weapon rather than a tool (p.26).

Additionally, the Company has provided its first [Responsible AI Transparency Report](#), which addresses: (i) how the Company builds generative applications responsibly; (ii) how it makes decisions about releasing generative applications; (iii) how it supports its customers in building responsibly; and (iv) how it learns, evolves, and grows, which includes a discussion of how industry and others stand to significantly benefit from the key role that governments can play regarding consensus-based safety frameworks (pp.4-32).

Regarding oversight of this issue, the [environmental, social, and public policy committee](#) assists the board in overseeing the key non-financial regulatory risks that may have a material impact on the Company and especially its ability to sustain trust with customers, employees, and the public, which includes policies and programs and related risks that concern environmental sustainability, the social and public policy impacts of technology including privacy, digital safety, and responsible artificial intelligence, and legal, regulatory, and compliance matters relating to competition / antitrust, trade, and national security. The committee also reviews and provides guidance to the board and management about key environmental and social matters such as responsible artificial intelligence and human rights, among others.

Summary

Analyst Note

While the Company provides significant and meaningful disclosure regarding its human rights efforts, it does not appear to provide details specifically concerning its military contracts. The Company's HRIA published in 2023 also states that stakeholders want more disclosure regarding the Company's military contracts, and they feel there is a lack of transparency regarding the Company's product design and deployment.

RECOMMENDATION

We recognize that there could be significant reputational and financial risks stemming from the Company's involvement in the development of weapons used by the military for training and/or combat purposes. We also recognize that the proponent appears to largely take issue with how government agencies are using the Company's technology, as opposed to actions undertaken directly by the Company. While there is a need for companies to monitor the reputational impacts of how their customers are using its products, we also appreciate the difficulty in the Company policing and assessing how government agencies are utilizing its technology. However, given these issues could present material risks for the Company, we believe that additional scrutiny on this matter is warranted.

We note that, in its statement of opposition, the Company clearly outlines its position on providing technology to the military. However, we believe that additional disclosure on this matter could be warranted. Moreover, this sentiment is highlighted by the Company's 2023 third-party HRIA, which made a specific recommendation for the Company to "consider conducting theme-specific due diligence exercises or additional HRIAs focused on," among other things, "assessing the company's commercial relationships with military agencies and their impacts on BIPOC and other vulnerable communities" and "other human rights challenges prioritized by the stakeholder organizations with which it engages." In the past year, the Company does not appear to have provided any further disclosure related to its commercial relationships with military agencies regarding human rights impacts.

To be clear, we do not believe that the Company should necessarily be limited in what contracts it can pursue or with which customers it may conduct business. However, given the financial, reputational and human-capital-related risks that could be realized via its work with the military, we believe that a more robust discussion on this issue should be disclosed. While we are not of the view that a third party would need to be engaged in order to conduct this report, as requested by this resolution, we do believe that more robust disclosure concerning these matters is warranted at this time. We believe that such disclosure would provide shareholders with a more meaningful understanding of these considerations or more fully allow them to evaluate how the Company is managing attendant risks. As such, we believe that support for this proposal is warranted at this time.

We recommend that shareholders vote **FOR** this proposal.

5.00: SHAREHOLDER PROPOSAL REGARDING ASSESSMENT OF INVESTMENTS IN BITCOIN

AGAINST

PROPOSAL REQUEST:	That the board conduct an assessment to determine if including Bitcoin in the Company's balance sheet is in shareholders' best long-term interests	SHAREHOLDER PROPONENT:	National Center for Public Policy Research
BINDING/ADVISORY:	Precatory	REQUIRED TO APPROVE:	Majority of votes cast
PRIOR YEAR VOTE RESULT (FOR):	N/A		
RECOMMENDATIONS, CONCERNS & SUMMARY OF REASONING:			
AGAINST - Not in the best interests of shareholders			

GLASS LEWIS REASONING

- We believe that management and the board typically have more and better information about the Company and its operations and are, therefore, in the best position to make decisions regarding the Company's investment strategies, absent egregious behavior or disregard of shareholder interests.

PROPOSAL SUMMARY

Text of Resolution: *Resolved: Shareholders request that the Board conduct an assessment to determine if diversifying the Company's balance sheet by including Bitcoin is in the best long-term interests of shareholders.*

Proponent's Perspective

- Corporations have a fiduciary duty to maximize shareholder value not only by working to increase profits but also by working to protect those profits from debasement;
- The proponent expresses concern regarding the average inflation rate in the U.S. over the last four years, stating the true inflation rate is significantly higher than CPI measures;
- As of March 31, 2024, the Company had \$484 billion in total assets, the plurality of which are U.S. government securities and corporate bonds that barely outpace inflation;
- As of June 25, 2024, the price of Bitcoin increased by 99.7% over the previous year, outperforming corporate bonds by roughly 94% on average, and over the past five years, the price of Bitcoin increased by 414%, outperforming corporate bonds by roughly 411% on average;
- The institutional and corporate adoption of Bitcoin is becoming more commonplace;
- The stock of another technology company that holds Bitcoin on its balance sheet has significantly outperformed the Company's stock this year, despite doing only a fraction of the business that the Company has;
- The Company's second largest shareholder, BlackRock, offers its clients a Bitcoin ETF;
- Bitcoin is a more volatile asset, at the moment, than corporate bonds, so companies should not risk shareholder value by holding too much of it, but companies should also not risk shareholder value by ignoring Bitcoin altogether; and
- At minimum, companies should evaluate the benefits of holding some, even just 1%, of their assets in Bitcoin.

Board's Perspective

- The Company's management already carefully considers the topic of this proposal;
- The Company's Global Treasury and Investment Services team evaluates a wide range of investable assets to fund the Company's ongoing operations, including assets expected to provide diversification and inflation protection, and to mitigate the risk of significant economic loss from rising interest rates;
- Past evaluations have included Bitcoin and other cryptocurrencies among the options considered, and the Company continues to monitor trends and developments related to cryptocurrencies to inform future decision-making;
- As this proposal itself notes, volatility is a factor to consider in evaluating cryptocurrency investments for corporate treasury applications that require stable and predictable investments to ensure liquidity and operational funding; and
- The Company has strong and appropriate processes in place to manage and diversify its corporate treasury for the long-term benefit of shareholders, and this requested public assessment is unwarranted.

THE PROPONENT

The National Center for Public Policy Research

The proponent of this proposal is the National Center for Public Policy Research ("NCPPr"). The NCPPr describes itself as a "communications and research foundation supportive of a strong national defense and dedicated to providing free-market solutions to today's public policy problems" that believes "the principles of a free market, individual liberty and personal responsibility provide the greatest hope for meeting the challenges facing America in the 21st century." As NCPPr is not an investor, it does not appear to have AUM. The NCPPr [states](#) that it "trains its sights" on the following issues: (i) shareholder activism; (ii) new leadership for Black America (Project 21) (iii) environmental policy; (iv) regulatory policy; (v) fiscal policy, health care, and retirement security; (vi) government accountability and legal reform; (vii) national defense; (viii) national sovereignty; and (ix) emerging issues.

A division of the NCPPR, the [Free Enterprise Project](#) ("FEP") [describes](#) itself as the "original and premier opponent of the woke takeover of American corporate life and defender of true capitalism." It files shareholder resolutions, engages corporate CEOs and board members, submits public comments, engages state and federal leaders, crafts legislation, files lawsuits and directs media campaigns to push corporations to respect their fiduciary obligations and to stay out of political and social engineering." Based on the disclosure provided by companies concerning the identity of proponents, during the first half of 2023, NCPPR submitted 32 shareholder proposals to a vote, receiving an average of 1.8% support (with none receiving majority shareholder support).

GLASS LEWIS ANALYSIS

Glass Lewis recommends that shareholders take a close look at proposals such as this to determine whether the actions requested of the Company will clearly lead to the enhancement or protection of shareholder value. Glass Lewis believes that directors who are conscientiously exercising their fiduciary duties will typically have more and better information about the Company and its situation than shareholders. Those directors are also charged with making business decisions and overseeing management. Our default view, therefore, is that the board and management, absent a suspicion of illegal or unethical conduct, will make decisions that are in the best interests of shareholders. However, while the board and management are responsible for making these business decisions, we also believe that it is important that shareholders be provided with sufficient information to allow them to gauge the effectiveness and rigor of the management and oversight of key issues.

BACKGROUND

This proposal requests that the Company undertake an assessment to determine whether including Bitcoin in the Company's balance sheets is in the best long-term interests of shareholders. [Bitcoin](#) was the world's first digital currency, conceptualized in the early 1980s but not launched as a usable digital currency until 2009, when Satoshi Nakamoto, a pseudonym whose true identity has never been verified, launched Bitcoin in January of that year. Nakamoto also introduced the blockchain system, the backbone of the cryptocurrency market. [Cryptocurrency](#) is made usable through blockchain technology, which is a digital ledger of transactions replicated and distributed across a network of computer systems, securing the information. Though some countries (such as China, India, and Saudi Arabia) have banned cryptocurrency mining and trading, many other countries have fully embraced it. For example, the U.S. Securities and Exchange Commission ("SEC") approved 11 spot Bitcoin exchange-traded funds in January 2024, and in July 2024 gave final approval for spot Ether ETFs to start trading, adding legitimacy to the asset class (Julie Pinkerton. "[The History of Bitcoin](#)." *U.S. News and World Report*. October 23, 2024).

The two most widely used cryptocurrencies are Bitcoin and Ether. When combined, both currencies account for about 70% of the cryptocurrency global market cap and more than half of all trading volume as of May 31, 2024 (Amy C. Arnott, CFA. "[How to Use Bitcoin in Your Portfolio](#)." Morningstar. June 4, 2024). Despite its growing popularity, Bitcoin has exhibited notable volatility since it entered the markets in 2009, an issue noted by both the Company and the proponent. While it was being championed as the best investment of the year in 2013 by *Forbes*, a year later, *Bloomberg* countered that it was the worst investment in 2014. Further, controversy ensued when FTX, the leading cryptocurrency exchange by trading volume, declared bankruptcy in late 2022 as panicked customers created an \$8 billion liquidity shortfall (Julie Pinkerton. "[The History of Bitcoin](#)." *U.S. News and World Report*. October 23, 2024).

An important feature of Bitcoin, unlike traditional currencies such as the dollar or pound, is that it is not controlled by centralized financial institutions. While this feature is appealing to some investors, it leaves the digital currency vulnerable to extreme volatility. Another unique feature of Bitcoin is that there is not an infinite supply of bitcoins, unlike other digital currencies. For example, the amount of Bitcoin that can be mined by "miners" (who validate the blockchain transactions and are then paid in cryptocurrency) is capped at 21 million, with most bitcoins already in circulation. As a result, approximately every four years, once the Bitcoin blockchain achieves a certain size, the amount rewarded to "miners" decreases by half. Although this extends the supply of Bitcoin, not all investors may be satisfied with smaller returns for mining, especially given the significant amount of energy required to update the blockchain (Brandon Drenon, Joe Tidy, Liv McMahon. "[What is Bitcoin? Key Cryptocurrency Terms and What They Mean](#)." *BBC*. April 23, 2024).

Asset management firms like BlackRock have [explored](#) Bitcoin as a unique diversifier. BlackRock and Fidelity Investments both launched Bitcoin ETFs on January 11, 2024, representing a watershed moment for cryptocurrency. BlackRock's [iShares Bitcoin Trust ETF](#) and the [Fidelity Bitcoin ETF](#) both make Bitcoin more accessible to investors (Katie Greifeld, Sidhartha Shukla. "[BlackRock's \\$20 Billion ETF Is Now the World's Largest Bitcoin Fund](#)." *Bloomberg*. May 29, 2024). According to BlackRock, Bitcoin [appeals](#) to investors due to its detachment from traditional risk and return drivers.

However, according to a February 2024 [survey](#) from the Pew Research Center, 63% of Americans have little to no confidence that current methods of investing in, trading, or using cryptocurrency are safe and reliable. It further noted that the number of Americans who have used cryptocurrency has not grown in the past three years. Since 2021, only 17% of

U.S. adults say they have ever invested in, traded, or used a cryptocurrency. Additionally, among those who are familiar with but have not invested in cryptocurrency, 82% say they are not very or not at all confident in it, compared to 39% among those who have invested in cryptocurrency.

Without wider acceptance of Bitcoin as a currency, some contend that its value appears to mostly hinge on market sentiment and speculation, rather than intrinsic value, which makes it a risky asset. As such, some have argued that Bitcoin's frequent price fluctuations have undermined its reliability as a store of value or as a hedge against inflation, at least in the short-term (Mike Winters. " [Bitcoin's Dizzying Price Movements Make it a Risky Investment, Say Investing Experts: 'It's Pure, Unadulterated Speculation'](#)." *CNBC*. August 8, 2024). Morningstar recommends holding cryptocurrency for at least ten years, based on how long it usually takes to recover from a drawdown. Morningstar also notes that Bitcoin, Ether, and other cryptocurrencies have become less valuable as portfolio diversifiers because their correlations with other major asset classes have steadily risen in recent years. As a result, it has found that there is no guarantee that adding a cryptocurrency will improve a portfolio's risk-adjusted returns (Amy C. Arnott, CFA. "[How to Use Bitcoin in Your Portfolio](#)." Morningstar. June 4, 2024).

RECOMMENDATION

Glass Lewis believes that a well-functioning, informed board of directors should receive reasonable deference (though not complete deference) from shareholders on matters such as its investment strategies. Such a board is often in the best position, with more information and experts at its disposal, to assess a company's investment needs in relation to its operational plans, growth prospects and strategic alternatives. On proposals such as the present one, which ask shareholders to assert their judgment in place of the judgment of the board, we believe the burden is on the shareholder proponents to clearly demonstrate that the directors' judgment is incorrect and that the proposals, despite management opposition, will yield an increase in shareholder value. In the absence of some credible indication that the board has exercised poor judgment in its company oversight functions, we do not believe shareholders should substitute their judgment regarding the Company's investment strategies. Rather, shareholders should use their influence to push for a governance structure that protects shareholders and can hold directors accountable through a vote on the election of directors.

Moreover, it appears that, to a large extent, the Company has already fulfilled the request of this proposal, which requests that the board "conduct an assessment to determine if diversifying the Company's balance sheet by including Bitcoin is in the best interests of shareholders." Specifically, in response to this resolution, the Company states that:

Microsoft's Global Treasury and Investment Services team evaluates a wide range of investable assets to fund Microsoft's ongoing operations, including assets expected to provide diversification and inflation protection, and to mitigate the risk of significant economic loss from rising interest rates. Past evaluations have included Bitcoin and other cryptocurrencies among the options considered, and Microsoft continues to monitor trends and developments related to cryptocurrencies to inform future decision making.

Microsoft has strong and appropriate processes in place to manage and diversify its corporate treasury for the long-term benefit of shareholders and this requested public assessment is unwarranted.

(2024 DEF 14A, p.79)

Particularly in light of this assessment, we have no reason to suspect that the Company would not make investments in Bitcoin should it determine that it was in the best long-term interests of shareholders. As such, we do not find a clear showing by the proponents that shareholders should, in this instance, supplant the judgment of the board and management team or that adoption of this proposal will clearly lead to an increase in shareholder value. Accordingly, we are not convinced that adoption of this proposal is warranted at this time

We recommend that shareholders vote **AGAINST** this proposal.

6.00: SHAREHOLDER PROPOSAL REGARDING REPORT ON SITING IN COUNTRIES OF SIGNIFICANT HUMAN RIGHTS CONCERN

AGAINST

PROPOSAL REQUEST:	That the board commission a report assessing the siting of cloud data centers in countries of significant human rights concern	SHAREHOLDER PROPONENT:	Olga Bell Greenbaum D'Angelo and a co-filer
BINDING/ADVISORY:	Precatory		
PRIOR YEAR VOTE RESULT (FOR):	33.6%	REQUIRED TO APPROVE:	Majority of votes cast
RECOMMENDATIONS, CONCERNS & SUMMARY OF REASONING:			
AGAINST - Not in the best interests of shareholders			

SASB MATERIALITY	PRIMARY SASB INDUSTRY: Software & IT Services
	FINANCIALLY MATERIAL TOPICS: • Environmental Footprint of Hardware Infrastructure • Recruiting & Managing a Global, Diverse & Skilled Workforce • Managing Systemic Risks from Technology Disruptions • Data Privacy & Freedom of Expression • Data Security • Intellectual Property Protection & Competitive Behavior

GLASS LEWIS REASONING

- At this time, we find the Company's current human rights-related disclosure to be adequate with regard to the matters presented by the proponent and do not believe the proponent has provided compelling evidence that the Company's current policies, procedures, or practices with respect to the siting of its cloud data centers represent an imminent threat to shareholder value.

PROPOSAL SUMMARY

Text of Resolution: *Resolved: Shareholders request the Board of Directors commission a report assessing the implications of siting Microsoft cloud datacenters in countries of significant human rights concern, and the Company's strategies for mitigating these impacts.*

The report, prepared at reasonable cost and omitting confidential and proprietary information, should be published on the Company's website within a year of the 2024 shareholders meeting.

Proponent's Perspective

- Shareholders are concerned by the Company's announced plans to expand datacenter operations to locations identified by the U.S. State Department's Country Reports on Human Rights Practices as presenting significant human rights challenges;
- Of particular concern is the plan to locate a Company datacenter in Saudi Arabia, as the U.S. State Department details the highly restrictive Saudi control of all internet activities and notes pervasive government surveillance, arrest, and prosecution of online activity;
- A former Twitter employee was recently convicted in federal court on charges of spying for Saudi Arabia, and reports suggest that Saudi authorities have recruited multiple spies inside U.S. Twitter operations to extract personal information and spy on private communications of exiled Saudi activists;
- The Company stated that the installation in Saudi Arabia would be consistent with its commitment to protecting fundamental rights and that its commitment to the Trusted Cloud Principles is an assurance of human rights protection, but it has refused to disclose how these assurances will be enacted or enforced, as there is no recognized oversight body for the principles; and
- The Company has provided no evidence that it has conducted a human rights impact assessment or that it has engaged impacted stakeholders as required under the UN Guiding Principles on

Board's Perspective

- The proposal is unnecessary because the Company already discloses its human rights commitments and due diligence processes, and it provides ongoing public reporting on human rights;
- The Company seeks to operate responsibly and in accordance with its global commitment to respect human rights and the rule of law;
- When considering expansion of its datacenter footprint to new countries, the Company is guided by its Global Human Rights Statement, which requires that it conduct human rights due diligence in line with the UN Guiding Principles on Business and Human Rights;
- The Company's commitment to due diligence is global and ongoing, and it includes corporate-wide assessments, supply chain due diligence, and due diligence to inform operations in specific markets;
- To inform responsible market entry, including datacenter siting, the Company conducts extensive reviews, which include public assessment materials from trusted, independent third-party resources, and it also regularly sources additional input from outside counsel with international expertise in privacy and other human rights-related matters;
- As appropriate, in certain geographies, the Company conducts

Business and Human Rights ("UNGPs"), nor has the Company provided any disclosure of an assessment or mitigation plans.

heightened due diligence via country-level human rights impact assessments, and it uses what it learns from these studies to inform decision-making, develop and refine its policies and practices, mitigate risks, and improve its technologies and how it provides them to fulfill its commitment to human rights, updating these analyses as conditions and/or the Company's engagement in the geography changes;

- The Company followed its due diligence and risk mitigation process in evaluating the establishment of a new cloud datacenter region in Saudi Arabia and in determining it could be operated in a way consistent with the Company's commitment to protecting fundamental rights and focus on responsible cloud practices;
- Datacenter risk mitigation measures typically take the form of exclusion of specific types of services, technologies, and/or specific types of customers; and
- The Company worked with other major cloud service providers to develop, agree to, and publish the Trusted Cloud Principles.

GLASS LEWIS ANALYSIS

Glass Lewis recommends that shareholders take a close look at proposals such as this to determine whether the actions requested of the Company will clearly lead to the enhancement or protection of shareholder value. Glass Lewis believes that directors who are conscientiously exercising their fiduciary duties will typically have more and better information about the Company and its situation than shareholders. Those directors are also charged with making business decisions and overseeing management. Our default view, therefore, is that the board and management, absent a suspicion of illegal or unethical conduct, will make decisions that are in the best interests of shareholders.

We believe that it is prudent for firms to actively evaluate risks to shareholder value stemming from global activities and human rights practices along entire supply chains. This evaluation is particularly important for those companies with operations in sensitive regions prone to allegations of human rights abuses. As has been seen with other companies, allegations of human rights abuses can inflict, at a minimum, reputational damage on targeted companies and may dramatically affect shareholder value.

HUMAN RIGHTS CONVENTIONS

Universal Declaration of Human Rights

The notion of universally recognized human rights, applicable in both national and international jurisdictions, continues to evolve. Some important agreements that have been broadly ratified by the international community include the [Universal Declaration of Human Rights](#) ("UDHR"), the International Covenant on Civil and Political Rights, the Convention on the Elimination of All Forms of Racial Discrimination, the United Nations Convention Against Torture, the Convention of the Rights of the Child, and the International Convention on the Protection of the Rights of All Migrant Workers and Members of Their Families. Additionally, the [Geneva Conventions](#), which have been ratified by [196 countries](#), aim to protect the human rights of individuals involved in armed conflict. Further, in August 2003, the UN Sub-Commission on the Promotion and Protection of Human Rights published the "[Norms on the responsibilities of transnational corporations and other business enterprises with regard to human rights](#)." These norms propose to hold companies, as organs of society, responsible for promoting and securing human rights; however, they are currently neither binding nor monitored. In addition, in 2011, the UN Human Rights Council approved the [United Nations Guiding Principles on Business and Human Rights](#) ("Ruggie Principles") which state that "business enterprises should carry out human rights due diligence" including "assessing actual and potential human rights impacts, integrating and acting upon the findings, tracking responses, and communicating how impacts are addressed" (p.17).

COMPANY OPERATIONS IN THE MIDDLE EAST

The proponent cites concern regarding the Company's planned operations in Saudi Arabia, which has been cited by the U.S. State Department for having present significant human rights violations. The State Department [provides](#) annual Country Reports on Human Rights Practices, which cover internationally recognized individual, civil, political, and worker rights, as set forth in the Universal Declaration of Human Rights and other international agreements.

The Middle East

As mentioned by the proponent, although the Company does not currently have operations in regions in the Middle East, it announced in February 2023 that it intends to invest in a new cloud data center region in [Saudi Arabia](#). On November 1, 2023, the Company [issued](#) a news release discussing a visit by chair and CEO Satya Nadella to Saudi Arabia to experience first-hand how the latest advancements in cloud and AI are supporting the Kingdom's technology ecosystem of developers, startups, and organizations across every industry. While speaking at the "Microsoft AI, a New Era" event, Nadella met with local business leaders, government officials and developers and emphasized the role of AI in unlocking

new opportunities to accelerate the Kingdom's digital economy and transform the lives of its people.

In early January 2024, the Company said in an emailed statement to *Bloomberg* that it has a number of headquarters in the CEMA (Central & Eastern Europe, Middle East, & Africa) region, including one in Saudi Arabia. The Company was among several tech companies, including Amazon and Alphabet, that were reportedly building up their presence in Saudi Arabia in response to pressure from the government, which stated that it would no longer provide contracts to companies without regional headquarters in the country. The three U.S. companies, as well as others, each received licenses for the establishment of regional headquarters in Riyadh ahead of the January 1 deadline set by the Saudi government (Matthew Martin, Fahad Abuljadayel. "[Amazon, Microsoft Boosting Saudi Offices Amid State Pressure](#)." *Bloomberg*. January 8, 2024).

Meanwhile, the European Centre for Democracy and Human Rights ("ECDHR"), a non-profit organization, [states](#) that the establishment of a cloud region in the Kingdom of Saudi Arabia is projected to generate approximately \$24 billion in new revenue for the Company over the coming years. The ECDHR further states that the Company acknowledges the human rights risks associated with operating in Saudi Arabia but believes it can leverage its presence to advocate for change.

A number of significant human rights concerns have occurred in the region. For example, Yemen's political crises [began](#) in 2011 as a result of protests over poverty, unemployment, and political corruption. In 2014, the Iranian-backed Houthi militia took over Yemen's capital and declared control of the country. The Yemeni government enlisted the support of several military coalitions, including Saudi Arabia and the United Arab Emirates ("UAE"), for protection from the Houthi militia, and with this assistance gained control of 80% of the land. Recently, however, the UAE began supporting southern separatists in their efforts to establish an independent South Yemen, risking its partnership with Saudi Arabia (Mareike Transfeld. "[The UAE Is Weakening Its Partnership With the Saudis in Yemen. Here's Why That Matters](#)." *The Washington Post*. August 28, 2019). The UAE reduced its military presence in the country, but armed and trained tens of thousands of separatist fighters, signaling that its influence remains in the region (Marwa Rashad, Mohammed Mukhashaf. "[Saudi Arabia Defends Yemen Government Against UAE-Backed Separatists](#)." *Reuters*. September 5, 2019). Both Saudi Arabia and the UAE have called for a halt to conflicting military operations in the country and stated their continued support for the country's government (Patrick Wintour. "[Saudi Arabia and UAE Attempt to Paper Over Yemen Cracks](#)." *The Guardian*. September 9, 2019).

According to the UN, the [majority](#) of the Yemeni population still lives in Houthi-controlled land and has suffered from displacement, food insecurity, malnutrition, and the world's fastest-growing cholera outbreak. Approximately 80% of the population, or more than 24 million people, requires assistance and protection, while 10 million individuals rely on food aid for survival (Patrick Wintour. "[Yemen Civil War: The Conflict Explained](#)." *The Guardian*. June 20, 2019).

Military escalation in early 2022 stretched beyond Yemen's borders into the UAE and Saudi Arabia and led various world leaders to [call](#) for an urgent end to the violence in the hopes of still brokering peace through UN-led diplomatic efforts. Members of the UN security council have discussed the growing vulnerability of the Yemeni people, who may soon face the potential loss of funding for various humanitarian programs providing aid in the area. As of February 2022, President Biden was considering renewing the use of a terror designation for Houthi rebels, due to their increased use of missile attacks, but aid groups fear it would further harm the people of Yemen (Missy Ryan, John Hudson. "[Missile Attacks Fuel Support for Reversing U.S. Stance and Placing Yemen Rebels Back on Terrorist Blacklist](#)." *Washington Post*. February 11, 2022). Further, the UN released a report in March 2022, warning of an extreme increase in the number of people in Yemen expected to experience famine this year due to the ongoing war, estimating that as many as 161,000 Yemenis will face hunger (Samy Magdy. "[Harrowing figures: Yemen Report Says 161K to Face Famine](#)." *AP News*. March 14, 2022).

Since entering the conflict in 2015, Saudi Arabia has become the world's largest weapons importer and the U.S. has become its biggest supplier (Julian Borger. "[US Leads Arms Sales to Saudis, Followed by UK, From 2014-2018](#)." *The Guardian*. October 3, 2019). Nearly one-fifth of U.S. weapons exports went to Saudi Arabia in the five years ending in 2017, making it the largest market for American defense contractors. Saudi Arabia has spent close to \$90 billion on weapons and defense systems from U.S. contractors since the 1950s, including approximately \$5.5 billion in 2017 alone (Rachel Layne. "[U.S. Defense Companies Dodge Flak on Saudi Arabia Weapons Sales](#)." *CBS News*. October 26, 2018). Despite the ongoing conflicts in Yemen and the murder of *Washington Post* journalist and dissident Saudi journalist Jamal Khashoggi, major defense manufacturers in the U.S. stated that they would stand by the Trump administration regarding continuing business with Saudi Arabia (Aaron Gregg, Christian Davenport. "[Defense Contractors Stand With White House on Saudi Arms Sales](#)." *The Washington Post*. October 25, 2018). In 2019, President Trump opted to bypass the long-standing precedent for congressional review of major weapons sales by declaring a national emergency to complete a sale of over \$8 billion worth of weapons to Saudi Arabia, the UAE, and Jordan (Patricia Zengerle. "[Defying Congress, Trump Sets \\$8 Billion-Plus in Weapons Sales to Saudi Arabia, UAE](#)." *Reuters*. May 24, 2019).

In 2021, the Biden administration reviewed sales to the Gulf States, in particular the agreement between Trump and the UAE (Michael Crowley. "[Biden Administration Reviewing Trump Arms Sales to U.A.E. and Saudi Arabia](#)." *The New York Times*. January 27, 2021). However, the Biden administration reportedly wants to keep a controversial Trump policy that

jump-started sales of armed drones to countries whose human rights records are under scrutiny. The Trump administration had reinterpreted a 33-year old agreement called the Missile Technology Control Regime ("MTCR") that allowed U.S. drone sales to governments in Taiwan, India, Morocco, and the UAE, which had previously been prohibited from buying them. The MTCR classified several of the most powerful U.S. drones as cruise missiles because they meet the technical specifications for unpiloted aircraft in the pact (Mike Stone. " [Exclusive: Biden Wants to Keep Trump Policy That Boosted Armed Drone Exports - Sources](#)." *Reuters*. March 25, 2021).

President Biden proposed another weapons sale to Saudi Arabia, which gained a bipartisan vote of confidence in December 2021, citing the country's need to defend itself from aerial cross-border attacks (Andrew Desiderio. " [Senate Backs Biden Admin Weapons Sale to Saudi Arabia](#)." *Politico*. December 7, 2021). In late 2021, the UAE announced that it was suspending talks with the U.S. regarding the F-35 sale after the deal had slowed due to concerns over how the UAE could operate the Lockheed Martin aircraft and how much of the aircraft's technology the country could access. However, the Secretary of State stated that the U.S. was prepared to move forward on the sale after a review of any technologies that are sold or transferred to partners in the region (Humeyra Pamuk, Rozanna Latiff. " [Blinken Says U.S. Ready to Move Forward With Sale of F-35s, Drones to UAE](#)." *Reuters*. December 15, 2021).

Despite that sale, Biden initiated a three-year ban on U.S. sales of offensive weapons to Saudi Arabia in 2021 in response to the country's campaign against Iran-aligned Houthis in Yemen. More recently, marking a shift in policy and a change in relations with the country, the State Department announced in August 2024 that it was lifting the suspension on certain transfers of air-to-ground munitions to Saudi Arabia. The State Department explained that it would examine new transfers on a case-by-base basis, consistent with the Conventional Arms Transfer Policy. As such, the Biden administration was expected to move ahead with a sale to Saudi Arabia as early as the week after the ban was lifted (Humeyra Pamuk, Patricia Zengerle, Steve Holland. " [US to Lift Ban on Offensive Weapons Sales to Saudi Arabia](#)." *Reuters*. August 9, 2024). By mid-October 2024, the Biden administration had approved the sale of billions of dollars in weapons to Saudi Arabia and the United Arab Emirates (Courtney McBride. " [US to Sell Up to \\$2.2 Billion in Weapons to UAE, Saudi Arabia](#)." *Bloomberg*. October 11, 2024).

COMPANY POLICIES AND DISCLOSURE

The Company [discusses](#) its data center operations in countries or regions with human rights challenges. It states that, when considering the expansion of its data center footprint to new countries, it is guided by its Global Human Rights Statement, which requires that it conduct due diligence in line with the UN Guiding Principles on Business and Human Rights ("UNGPs"). It adds that to inform responsible market entry, including datacenter siting, the Company conducts extensive reviews, which include public assessment materials from trusted, independent third-party resources like Freedom House's reports, the World Justice Project Rule of Law index, and the Transparency International Corruption Perceptions Index. Further, the Company regularly sources additional input from outside counsel with international expertise in privacy and other human rights-related matters. It also explains that as appropriate, in certain geographies, the Company conducts heightened due diligence via country-level human rights impact assessments, and it uses what it learns from these studies to inform decision-making, develop and refine its policies and practices, mitigate risks, and improve its technologies and how it provides them to fulfill its commitment to human rights. Moreover, the Company updates these analyses as conditions and/or the Company's engagement in the geography changes.

It continues to [state](#) that its data center risk mitigation measures typically take any or all of three of the following forms: (i) exclusion of specific types of services, such as consumer services, (ii) technologies, such as facial recognition technology, and/or (iii) specific types of customers, such as law enforcement agencies. The Company explains that in certain locations, it serves only enterprise customers, and does not store consumer data, to mitigate potential risks to the right to privacy and security, or it stores data outside of the jurisdiction, rather than on site. It clarifies that all uses of its online services must comply with the Company's Acceptable Use Policy, which prohibits use of the Company's services to violate the rights of others, among other restrictions, and it has the right to suspend users for violations.

The Company [adds](#) that, because customer and technical demand for more cloud computing infrastructure will continue to grow and it is not the only provider of such services, it worked with other major cloud service providers to develop, agree to, and publish the Trusted Cloud Principles. The Company states that, with these principles, it and other leading cloud computing services companies publicly commit to:

- Maintain consistent human rights standards as the Company competes for cloud service business across the globe;
- Work with governments to ensure the free flow of data, to promote public safety, and to protect privacy and data security in the cloud;
- Support the adoption of laws that allow governments to request data through a transparent process that abides by internationally recognized rule of law and human rights standards; and
- Work with the tech sector, public interest groups, and policymakers around the world, particularly in the countries where they operate or plan to operate data centers and cloud infrastructure, to ensure that laws and policies are substantially in line with the Trusted Cloud Principles.

With regard to its work with tech sector public interest groups and policymakers, the Company [affirms](#) that as part of its datacenter expansion in new countries, the Company reviews and assesses these laws and policies, and where necessary, it negotiates with governments to establish processes and conditions, grounded in its principles, under which it will respond to law enforcement requests for data.

In addition to the Company, other [signatories](#) to the Trusted Cloud Principles include Amazon, Google, Cisco, IBM, and Salesforce, among others.

The Company provides its [Global Human Rights Statement](#), which outlines its foundational principles and commitment to ongoing human rights due diligence. It states that its global and ongoing processes begin with a focus on identifying and assessing any actual, or potential, adverse human rights impacts that it may cause, contribute to, or be directly linked with, either through its own activities or as a result of its business relationships. It adds that its processes follow the UNGPs and the OECD Guidelines for Multinational Enterprises and that it conducts human rights impact assessments ("HRIAs"), to identify and prioritize salient risks. The Company states that it has conducted HRIAs at both the corporate and product levels and for various countries and locations and that its HRIA work includes regular engagement and consultation with stakeholders in an effort to understand and address the perspectives of vulnerable groups or populations. The Company also details its commitment to remediation and describes its channels to raise grievances or seek remediation regarding its human rights performance. Further, it discusses the Company's additional human rights commitments, including its commitment to support good governance and the rule of law, its commitment to engagement with people and governments in countries with significant human rights challenges, its commitment to vulnerable groups, and its commitment to human rights defenders. Additionally, the Company discusses key areas of impact, as well as internal governance.

The Company's 2024 [Impact Summary](#) discusses expanding its cyber skilling programs to 38 countries, partnering with nonprofits and educational institutions to train a new generation of diverse cybersecurity professionals (p.6). It also discusses how the Company protects fundamental rights, such as by promoting responsible business practices, expanding accessibility and connectivity, and advancing fair and inclusive societies (p.13). Throughout the impact summary, the Company further lists its goals and outlines its impact for each area.

Additionally, the Company provides details on its [Azure global infrastructure](#) and [datacenters](#), even [discussing](#) what a datacenter is and featuring information on datacenter communities. The Company [provides](#) transparency reports regarding a variety of topics, as well as providing a [Human Rights Annual Report](#) and [Microsoft Response: Business and Human Rights Resource Center](#), among others things.

Regarding oversight of this issue, the [environmental, social, and public policy committee](#) reviews and provides guidance to the board and management about key environmental and social matters, including human rights, and the Company's government relations activities.

Summary

Analyst Note

The Company maintains adequate human rights-related disclosure and policies, including those concerning the regions in which it operates. It also maintains board-level oversight of this issue.

RECOMMENDATION

Glass Lewis generally believes that a best effort should be put forward to identify and prevent risks to shareholder value stemming from human rights violations along its entire value chain. However, at this time, we find the Company's current disclosure to be adequate for shareholders to understand such risks. Further, we believe the proponent has failed to provide sufficiently compelling evidence that the Company's current policies, procedures, or practice with respect to human rights considerations represent an imminent threat to shareholder value.

We recognize that the Company could be exposed to significant legal, reputational, and regulatory risks. Accordingly, we believe that the Company should continue to provide shareholders with robust disclosure concerning how it is mitigating said risks. However, at this time, we are unconvinced that adoption of this proposal would add meaningfully to shareholders' understanding of how the Company was managing human rights-related issues at the various locations in which it maintains operations or that the Company has acted without regard to shareholder value in managing human rights-related risks. Thus, we do not believe support for this resolution is warranted at this time.

We recommend that shareholders vote **AGAINST** this proposal.

7.00: SHAREHOLDER PROPOSAL REGARDING REPORT ON RISKS OF PROVIDING AI TO FACILITATE NEW OIL AND GAS DEVELOPMENT AND PRODUCTION

AGAINST

PROPOSAL REQUEST:	That the Company report on the risks of providing advanced technology, including AI and ML tools to facilitate oil and gas development and production	SHAREHOLDER PROPONENT:	As You Sow and co-filers
BINDING/ADVISORY:	Precatory	REQUIRED TO APPROVE:	Majority of votes cast
PRIOR YEAR VOTE RESULT (FOR):	N/A		
RECOMMENDATIONS, CONCERNS & SUMMARY OF REASONING:			
AGAINST - Not in the best interests of shareholders			

SASB MATERIALITY

PRIMARY SASB INDUSTRY: Software & IT Services

FINANCIALLY MATERIAL TOPICS:

- Environmental Footprint of Hardware Infrastructure
- Recruiting & Managing a Global, Diverse & Skilled Workforce
- Managing Systemic Risks from Technology Disruptions
- Data Privacy & Freedom of Expression
- Data Security
- Intellectual Property Protection & Competitive Behavior

GLASS LEWIS REASONING

- Particularly in light of the Company's already robust disclosures, it is unclear that the requested reporting would materially improve shareholders understanding of how the Company is managing risks related to its provision of AI and machine learning tools to facilitate new oil and gas development and production.

PROPOSAL SUMMARY

Text of Resolution: *RESOLVED: Shareholders request that Microsoft report on the risks to the Company of providing advanced technology, including artificial intelligence and machine learning tools, to facilitate new oil and gas development and production.*

Supporting Statement: The report should address, at board discretion: the Company's exposure to financial, reputational, and competitive risks, including harm to employee hiring and retention, associated with the Company's deployment of advanced technology to facilitate new oil and gas development and production.

Proponent's Perspective

- Decarbonization of the energy system is required to stave off the most catastrophic consequences of climate change;
- The Company is widely seen as a global leader in reducing its contribution to climate change, but despite its stated commitments, the Company develops advanced technology platforms (including artificial intelligence ("AI"), machine learning, Internet of Things, cloud computing, and high performance computing (collectively, "advanced technology") that facilitate new oil and gas production projects that are likely to increase emissions and slow the transition to low-carbon technologies;
- The Company's advanced technologies enable continued fossil fuel expansion even though new, long-lead oil and gas projects have been found incompatible with the Paris Agreement's 1.5°C goal, and such projects and their associated GHG emissions are not factored into the Company's GHG emissions or reduction commitments;
- The Company's continued facilitation of new oil and gas extraction through advanced technology exposes it to material reputational, competitive, and operational risk;
- The development and deployment of advanced technology for the fossil fuel industry has caused outcry, and has led to staff resignations;
- A 2020 report assessing large technology companies found that

Board's Perspective

- The Company recognizes the world's need to both quickly develop access to more energy and to achieve an energy transition that achieves a net zero carbon economy by 2050, and the Company published a set of Energy Principles in 2022 specifically designed to guide the Company's work in addressing this challenge;
- The additional reporting requested in this proposal that would focus on a narrow customer segment is outside the approach of the global climate reporting standards the Company follows and is also unnecessary given the Company's existing disclosure of its approach to working with customers in the energy sector;
- The Company has worked with the energy industry for decades helping empower office workers and field operations across the entire value chain in mining, oil and gas, and power and utilities with a wide range of solutions and services;
- The Company has consistently demonstrated its commitment to sustainability and reducing its carbon footprint;
- While the Company acknowledges the concerns raised regarding the use of advanced technology in the fossil fuel industry, it is essential to recognize the broader context, given that technology plays a crucial role in helping various industries, including energy, transition towards more sustainable practices;
- The Company is committed to transparency in its operations and

the Company had the most contracts with oil and gas companies, while peers Google and Amazon had fewer ties, and in 2020, Google announced that it would no longer "build custom A.I. and machine learning algorithms to facilitate upstream extraction in the oil and gas industry";

- The Company has produced numerous reports on the responsible use of AI and AI-supported energy efficiency enhancements, but unlike Google, the Company continues to deploy advanced technologies for fossil fuel applications; and
- The Company's reporting on the climate-positive applications of advanced technologies, while ignoring the impact of new fossil fuel production, puts the Company at risk of accusations of greenwashing, which can have significant reputational or legal consequences.

reporting, and its existing disclosures are aligned with global multi-stakeholder frameworks and standards and provide shareholders with comprehensive information about its sustainability initiatives and the impact of its technologies;

- This proposal suggests that the Company's involvement with the fossil fuel industry exposes the Company to reputational and competitive risks, but the Company's balanced approach, which includes supporting the energy sector's transition to sustainable practices, mitigates these risks; and
- The Company values the input and engagement of its employees, and it has established channels for employees to voice their concerns and contribute to the Company's sustainability efforts.

The proponent has filed an [exempt solicitation](#) urging support for this proposal.

THE PROPONENT

As You Sow

As You Sow is a non-profit advocacy organization that "harness[es] shareholder power to create lasting change by protecting human rights, reducing toxic waste, and aligning investments with values." As You Sow is not an investor, and, therefore, does not have any assets under management, but uses investors' holdings to file [shareholder proposals](#) to "drive companies toward a sustainable future." It [states](#) that, since 1992, it has "utilized shareholder advocacy to increase corporate responsibility on a broad range of environmental and social issues." Areas of focus for As You Sow include [ocean plastics](#), [toxic chocolate](#), the climate and social impacts of [retirement funds](#), [climate change](#), [executive compensation](#), and [antibiotics and factory farms](#), among others.

Based on the disclosure provided by companies concerning the identity of proponents, during the first half of 2023, As You Sow submitted 32 shareholder proposals that received an average of 24.4% support, with one of these proposals receiving majority shareholder support.

GLASS LEWIS ANALYSIS

In general, we believe it is prudent for management to assess its potential exposure to all risks, including environmental and social issues and regulations pertaining thereto in order to incorporate this information into its overall business risk profile. When there is no evidence of egregious or illegal conduct that might suggest poor oversight or management of environmental or social issues that may threaten shareholder value, Glass Lewis believes that the management and reporting of environmental or social issues associated with business operations are generally best left to management and the directors who can be held accountable for failure to address relevant risks on these issues when they face re-election.

BACKGROUND

The oil industry is increasingly using AI and machine learning to drill faster, suggest better ways to frack, and predict when active well pumps will fail. The goal is for the new technology to reduce costs and help extract more oil from the ground. These uses ultimately help drillers and producers improve efficiency. Some expect significant cost savings, with 25% to 50% of cost savings in some scenarios as a result of using AI (David Wethe. ["AI Promises Faster Oil Drilling and Even More US Crude Supply."](#) *Bloomberg*. March 14, 2024). Further, the Company's technical architect, Azam Zaidi, highlighted in an April 2023 LinkedIn [blog post](#) that the Company's cloud service division was "enabling faster and more accurate decision-making and unlocking previously inaccessible reserves" in the fossil fuel industry (Maddie Stone. ["Microsoft Employees Spent Years Fighting the Tech Giant's Oil Ties. Now They're Speaking Out."](#) *Grist*. May 8, 2024).

This business line, however, raises concerns for the proponent. As part of its rationale for this proposal, the proponent states that the Company's continued facilitation of new oil and gas extraction through advanced technology exposes it to material reputational, competitive, and operational risk. In its exempt solicitation, the proponent [cites](#) international oilfield researcher Kimberlite, among other sources, to show that the Company has identified AI, machine learning, and cloud computing services (collectively "advanced technology") for the fossil fuel industry as a top growth opportunity. The cited [study](#) from Kimberlite indicates that the Company was the leading cloud provider serving the geological and geophysical community of the oil and gas industry, with Amazon in second place and Google in third place. The proponent also pointed to a 2020 [report](#) from Greenpeace that stated the Company appeared to have the most contracts with oil and gas companies, offering AI capabilities in all phases of oil production.

Whistleblower Allegations & SEC Complaint

The proponent also points to a September 2024 article from *The Atlantic*, which states that internal documents at the

Company and interviews with current and former staff indicate that, while the Company is trying to build a reputation as an AI leader for climate innovation, it is also marketing AI as a means to find and develop new oil and gas reserves and to maximize their production. According to the article, a whistleblower recently submitted internal Company documents to the U.S. Securities and Exchange Commission as part of a complaint alleging the Company has committed from public disclosures the "serious climate and environmental harms caused by the technology it provides to the fossil fuel industry." The whistleblower claims the information is of material and financial importance to investors (Karen Hao. "[Microsoft's Hypocrisy on AI](#)." *The Atlantic*. September 13, 2024).

Employee Resignations

Bolstering those allegations is a May 2024 report by *Grist* stating that some Company employees resigned in early 2024 over ethical concerns about the Company's work to automate and accelerate oil extraction for the fossil fuel industry. One former employee, Holly Alpine, who helped organize the Company's Sustainability Connected Community, stated that her resignation "was driven in part by the realization that the tech industry, including the Company, is increasing the profitability and competitiveness of these fossil fuel giants and perpetuating their existence when they should be phased out." The Company responded, highlighting its support for employees' roles in the Company's sustainability mission and adding:

The energy transition is complex and requires moving forward in a principled manner. We believe that technology has an important role to play in helping the industry decarbonize, and that requires balancing the energy needs and industry practices of today while inventing and deploying those of tomorrow.

And we continually monitor our emissions, accelerate progress while increasing our use of clean energy to power data centers, purchasing renewable energy and other efforts to meet our sustainability goals of being carbon negative, water positive, and zero-waste by 2030.

During a town hall meeting with employees in 2019, the Company's CEO responded to an employee's question about the ethical use of the Company's technology for extraction at fossil fuel companies. According to a meeting transcript, the CEO emphasized that fossil fuel companies were actively investing in the energy transition, implying that the Company was enabling them to put more resources into emissions-reducing innovations. However, Alpine noted that the Company continued to build partnerships based on boosting production for fossil fuel producers (Maddie Stone. "[Microsoft Employees Spent Years Fighting the Tech Giant's Oil Ties. Now They're Speaking Out](#)." *Grist*. May 8, 2024).

Employee Criticism

In 2021, employees participating in the Company's Sustainability Connected Community drafted a memo to the Company's leadership, criticizing the Company for "enabling far more emissions than [it offsets or removes]" by customizing its cloud computing technologies for fossil fuel companies. In the memo, employees estimated that a deal with Exxon Mobil to expand production in Texas and New Mexico by up to 50,000 barrels a day could enable carbon emissions equivalent to 640% of the Company's carbon removal target for 2021. The memo also outlined over a dozen recommendations for the Company, including ceasing to develop customized software tools for the purpose of increasing oil extraction. Company employees then participated in a meeting with executives about the memo and ways to respond to employee concerns (Maddie Stone. "[Microsoft Employees Spent Years Fighting the Tech Giant's Oil Ties. Now They're Speaking Out](#)." *Grist*. May 8, 2024).

In a blog post in March 2022, the Company discussed the principles that guide its work with the energy industry and suggested that it would only develop specialized tools for oil and gas extraction for companies that had agreed to reach net zero emissions by 2050 or sooner. However, employees became frustrated because the goal only addressed companies' net zero targets (covering Scope 1 and 2 emissions) and companies only had to provide a target with no evidence that they were on track to meet those targets. Employees subsequently held several meetings with Company executives regarding their frustrations over the Company's commitments and their conflicting work with the fossil fuel industry (Maddie Stone. "[Microsoft Employees Spent Years Fighting the Tech Giant's Oil Ties. Now They're Speaking Out](#)." *Grist*. May 8, 2024).

COMPANY AND PEER DISCLOSURE

Company Disclosure

In its response to this proposal, the Company states that:

While we acknowledge the concerns raised regarding the use of advanced technology in the fossil fuel industry, it is essential to recognize the broader context. Technology plays a crucial role in helping various industries, including energy, transition towards more sustainable practices. By working with energy companies, we can drive innovation and support the development of low- and zero-carbon solutions to accelerate the clean energy transition.

(2024 DEF 14A, p.83)

With respect to its climate commitment, the Company affirms in its latest [Environmental Sustainability Report](#) fact sheet that it announced in January 2020 that it would be carbon negative by 2030 and that by 2050 it would remove from the atmosphere an equivalent amount of all the carbon the Company has emitted either directly or by its electricity consumption since being founded in 1975. It specifies that the Company plans to achieve this goal by reducing Scope 3 emissions (market-based and management's criteria) by more than half, and by reducing Scope 1 and 2 (market-based) emissions to near zero by the middle of the decade through energy efficiency work and reaching 100% renewable energy by 2025 (p.8).

The Company released a [blog post](#) in September 2019 in which it announced the Company's collaboration with Schlumberger and Chevron to accelerate the digital transformation. The Company's CEO stated the partnership would apply the power of Azure to unlock new AI-driven insights to address energy challenges, including sustainability. Another 2019 [message](#) from the Company states that through a digital partnership with ExxonMobil, the firm would be able to expand production in the Permian Basin by as much as 50,000 oil-equivalent barrels a day by 2025. The Company stated that the project was the largest-ever oil and gas acreage to use cloud technology and that the Permian application would generate billions of dollars in value over the next decade. Similarly, in 2017, the Company [announced](#) a strategic alliance with Halliburton to drive digital transformation across the oil and gas industry in which Halliburton would use the Company's cloud services and machine learning, among other tools, to promote exploration and production.

Subsequently, the Company responded to a May 2020 Greenpeace report about its role in helping fossil fuel companies locate and extract oil and gas by releasing a [statement](#) which asserted:

We agree that the world confronts an urgent carbon problem and we all must do more and move faster to reach a net zero-carbon future. The reality is that the world's energy currently comes from fossil fuels and, as standards of living around the world improve, the world will require even more energy. That makes realizing a zero-carbon future one of the most complex transitions in human history.

We're up for the challenge. That's why we have committed to be carbon negative by 2030 and to removing all the carbon we've emitted since our founding by 2050. We will shift to 100% supply of renewable energy for the carbon-emitting electricity consumed by all our data centers, buildings and campuses by 2025. Technology can accelerate the transition to a zero-carbon future, so we are helping our customers reduce their carbon footprints and co-innovating low-carbon solutions. We're also using \$1 billion of our capital for climate innovation, including developing new technologies for carbon capture and removal. We are also supporting policies that advance an inclusive transition to a low-carbon future.

We're encouraged by the growing number of energy sector commitments to transitioning to cleaner energy and lowering carbon emissions, but they can't do it alone. Businesses, governments and civil society can rise to the challenge to meet the world's growing energy demands and achieve a net zero carbon future.

In March 2022, the Company [addressed](#) its evolving work with energy companies, shared an [update](#) on progress towards its sustainability goals, and provided its [Energy Principles](#), which state:

- The Company will increase its engineering investments and efforts in low- and zero-carbon energy businesses to help accelerate the energy transition;
- The Company is committed to helping all customers, including all energy customers, in the development of effective net zero commitments;
- The Company will sell its commercially available software technology and cloud services to all customers, inclusive of energy customers; and
- The Company may provide technical and engineering resources to develop or co-develop specialized services for subsurface exploration and extraction of fossil fuels with energy customers who have publicly committed to net zero carbon targets.

Regarding this last energy principle, the Company [specified](#) that the net zero carbon target cover Scope 1 and 2 emissions and be attained by 2050 or sooner. Additionally, the Company asserted that it was guided by a desire to help drive impact (to encourage invention and innovation) while not being prescriptive about proposed solutions.

The Company also discloses its [playbook](#) for accelerating sustainability with AI, and states that AI is a tool to help accelerate the deployment of existing sustainability solutions and the development of new ones -- faster, cheaper, and better (p.4). The Company then discusses three game-changing abilities of AI to help society overcome key bottlenecks to net zero progress, including:

- Measure, predict, and optimize complex systems;
- Accelerate the development of sustainability solutions; and
- Empower the sustainability workforce.

(p.5)

It [details](#) its five-point playbook for AI and sustainability:

1. Invest in AI to accelerate sustainability solutions;
2. Develop digital and data infrastructure for the inclusive use of AI for sustainability;
3. Minimize resource use in AI operations;
4. Advance AI policy principles and governance for sustainability; and
5. Build workforce capacity to use AI for sustainability.

(p.6). The playbook [links](#) to the Company's [carbon policy](#) and [electricity policy](#) briefs. While the carbon policy brief does not appear to address AI, the electricity policy brief states that advanced capabilities like AI can speed the pace of transition, help effectively integrate new technologies on the grid, and enhance the cybersecurity of critical power sector infrastructure. It adds that policies that promote a transparent, standardized approach to electricity data will help enhance operational performance and increase data availability and transparency. Further, it explains that this will allow for better accounting of real-time carbon emissions to show where the biggest carbon reduction potential lies and inform clear energy development decisions (p.6). The Company also [discusses](#) tracking AI's impact on the global race to net zero, and it explores in detail three questions:

1. How much energy is the global expansion of AI compute likely to consume?
2. How fast will the world's electric grids decarbonize?
3. To what extent will AI enable sustainability solutions?

(pp.21-22)

Regarding oversight of this issue, the [environmental, social, and public policy committee](#) assists the board in overseeing the key non-financial regulatory risks that may have a material impact on the Company and especially its ability to sustain trust with customers, employees, and the public, which includes policies and programs and related risks that concern environmental sustainability, the social and public policy impacts of technology including privacy, digital safety, and responsible artificial intelligence, and legal, regulatory, and compliance matters relating to competition / antitrust, trade, and national security. The committee also reviews and provides guidance to the board and management about key environmental and social matters such as climate change, environmental sustainability, responsible artificial intelligence, and human rights, among others.

Peer Disclosure

To compare, **Alphabet Inc.** (NASDAQ: GOOGL) also responded to the Greenpeace report in May 2020 by confirming that it was helping fossil fuel companies locate and extract oil and gas deposits, but that it would not "build custom AI/ML algorithms to facilitate upstream extraction in the oil and gas industry." A spokesperson for Google added that the firm took approximately \$65 million from oil and gas companies in 2019, which accounted for less than 1% of its total Google Cloud revenues, and given its relatively small portion of the oil and gas market, it was easy for the firm to vow not to compete using AI/ML in that industry (Sam Shead. "[Google Plans to Stop Making A.I. Tools for Oil and Gas Firms.](#)" *CNBC*. May 21, 2020).

Regarding board oversight, the [audit and compliance committee](#) has responsibility for oversight of risks and exposures associated with data privacy and security, competition, legal, regulatory, compliance, civil and human rights, sustainability, and reputational risks. Alphabet also states that the board's oversight function of major risks and risk exposures, including those relating to or resulting from the firm's development and implementation of AI in its products and services, sits at the top of its risk management framework, and as set forth in Alphabet's Corporate Governance Guidelines, the board is ultimately responsible for covering strategic, financial, and execution risks and exposures associated with the firm's business strategy, production innovation, and policy and significant regulatory matters that may present material risk to its financial performance, operations, plans, prospects, or reputation. Further, it adds that the board's skills and expertise, including deep technical expertise in computer science, facilitates oversight of a highly complex global business, and the full board meetings have regularly and extensively covered AI issues. The audit committee and senior management provide the board with reports and updates regarding issues and risk exposures regarding AI development, and these discussions ensure that the board is fully involved in the oversight of Alphabet's business strategies and plans as they relate to AI (2024 DEF 14A, p.94).

To further compare, **Meta Platforms, Inc.** (NASDAQ: META) discusses responsible AI and AI for climate in its most recent [Sustainability Report](#), stating that its climate-related AI initiatives, including [Open Catalyst](#) and the [Llama Impact Grants](#) program, are advancing climate action for the broader community. Meta affirms that it is responsibly advancing AI and sharing its data, models, and learnings with the AI and broader community. It also discusses developing an open source model for sustainable concrete and using AI to map the Earth's forests (pp.66-69). The firm states that it strives to have a positive impact on the communities and environment where it operates, and that its size and scale enable it to influence the future of decarbonization, the expansion of renewable energy, and the integrity of climate solutions (p.6).

Further, Meta discloses its [sustainability strategy](#), which is anchored by three components and defines how the firm operates. Its sustainability strategy protects the well-being of people and the planet by: (i) reducing Meta's environmental footprint, championing renewable energy, restoring water resources and engaging its suppliers; (ii) respecting human rights and protecting the environment within Meta's operations and supply chain; and (iii) boosting energy and water efficiency in its data centers. The firm also [provides](#) its commitment to responsible AI, with details on its five pillars of: (i) privacy and security; (ii) fairness and inclusion; (iii) robustness and safety; (iv) transparency and control; and (v) accountability and governance. It discusses how Meta is actioning responsible AI through datasets, privacy, ad delivery, associations, ad-driven feeds, system cards, and policy. However, Meta does not appear to explicitly address providing artificial intelligence and machine learning tools to facilitate new oil and gas development and production.

Regarding oversight of this issue, Meta states that the board, both directly and through its committees, is actively involved in overseeing risks associated with AI, including misinformation and disinformation (2024 DEF 14A, p.83). Further, the [audit and risk oversight committee](#) assists the board in overseeing the risk management of Meta and oversees certain of the firm's major risk exposures, provided that the board may, in its discretion, exercise direct oversight with respect to any such matters. The committee also reviews with management, at least annually, the firm's major ESG risk exposures and the steps management has taken to monitor or mitigate such exposures, in coordination with the other committees of the board as appropriate. **It also periodically reviews with management the status of the firm's ESG program and strategy.**

Summary	
GRI/SASB-Indicated Sustainability Disclosure	GRI (pp.16-19)
Peer Comparison	Google is the only one of the three companies to explicitly commit not to building custom AI/ML algorithms to facilitate upstream extraction in the oil and gas industry.
Analyst Note	As part of its Energy Principles, the Company may provide technical and engineering resources to develop or co-develop specialized services for subsurface exploration and extraction of fossil fuels with energy customers who have publicly committed to net zero carbon targets to be attained by 2050 or sooner. Additionally, the Company maintains explicit board-level oversight of climate change, environmental sustainability, and responsible artificial intelligence.

RECOMMENDATION

We understand the proponent's concerns regarding the Company's development of AI and machine learning for companies in the oil and gas industry. The proponent contends that its work with these companies could expose the Company to financial, reputational, and competitive risks, including harm to employee hiring and retention, associated with its deployment of advanced technology to facilitate new oil and gas development and production. However, we ultimately do not believe the proponent has sufficiently argued that the Company's management of this issue presents a financially material risk.

In our view, the proponent has not articulated any tangible financial harm to the Company on account of its actions with regard to its business with oil and gas companies. Further, the proponent notes widespread employee dissent and suggests that the Company has failed to attract and retain top talent. However, in support of this claim, the proponent only points to the departure of two employees on account of this issue. Although we understand that this could have wider impacts within the organization, it is not clear that such impacts are imminent or likely.

Moreover, we believe that the Company has provided significant disclosure regarding its sustainability commitments, including its own net zero commitment and its Energy Principles, which require a net zero target for Scope 1 and 2 emissions, achieved by 2050 or sooner, for any fossil fuel company to whom it would extend specialized services for subsurface exploration and extraction of fossil fuels. The Company also acknowledges the concerns raised regarding the use of advanced technology in the fossil fuel industry, and states that it is essential to recognize the broader context, that technology plays a crucial role in helping various industries, including energy, transition towards more sustainable practices. Given this disclosure, it is unclear that additional disclosure in this regard would materially improve shareholders understanding of how the Company is managing this matter.

Given the above, we are not convinced that support for this proposal is warranted at this time. However, we will continue to monitor how the Company is handling this issue and whether additional risks have materialized and may recommend in favor of this proposal at future AGMs should it be clear that the management of this issue presents a clear risk to shareholder value.

We recommend that shareholders vote **AGAINST** this proposal.

8.00: SHAREHOLDER PROPOSAL REGARDING REPORT ON AI MISINFORMATION AND DISINFORMATION

FOR

PROPOSAL REQUEST:	That the Company report on the risks presented by its role in facilitating AI misinformation and its plans to remediate those harms	SHAREHOLDER PROPONENT:	Arjuna Capital and co-filers
BINDING/ADVISORY:	Precatory		
PRIOR YEAR VOTE RESULT (FOR):	21.2%	REQUIRED TO APPROVE:	Majority of votes cast
RECOMMENDATIONS, CONCERNS & SUMMARY OF REASONING:			
FOR -	Information concerning exposure to risks related to misinformation and disinformation could be decision-useful for shareholders		

SASB MATERIALITY

PRIMARY SASB INDUSTRY: Software & IT Services

FINANCIALLY MATERIAL TOPICS:

- Environmental Footprint of Hardware Infrastructure
- Recruiting & Managing a Global, Diverse & Skilled Workforce
- Managing Systemic Risks from Technology Disruptions
- Data Privacy & Freedom of Expression
- Data Security
- Intellectual Property Protection & Competitive Behavior

GLASS LEWIS REASONING

- We believe that a discussion of the efficacy of the Company's policies aimed at preventing the dissemination of misinformation and disinformation generated via its AI technologies would allow shareholders more insight into how the Company is managing these potentially material risks

PROPOSAL SUMMARY

Text of Resolution: *Resolved, Shareholders request the Board issue a report, at reasonable cost, omitting proprietary or legally privileged information, to be published within one year of the Annual Meeting and updated annually thereafter, assessing the risks to the Company's operations and finances as well as risks to public welfare presented by the company's role in facilitating misinformation and disinformation disseminated or generated via artificial intelligence, and what steps, if any, the company plans to remediate those harms, and the effectiveness of such efforts.*

Proponent's Perspective

- There is widespread concern that generative Artificial Intelligence ("AI") may dramatically increase misinformation and disinformation globally, posing serious threats to democracy and democratic principles;
- The Company has reportedly invested over \$13 billion in OpenAI, the company that developed ChatGPT along with the Company, and has integrated ChatGPT into its AI-powered digital assistant Copilot;
- Generative AI's disinformation may pose serious risks to democracy by manipulating public opinion, undermining institutional trust, and swaying elections;
- Many legal experts believe technology companies' liability shield provided under Section 230 of the Communications Decency Act may not apply to content generated by AI;
- Shareholders are concerned that generative AI presents the Company with significant legal, financial, and reputational risk, and the Company has already faced substantial defamation litigation due to misinformation produced by the Company's generative AI;
- The Company will need to be responsive to the evolving AI regulatory landscape, including the EU's AI Act, Biden's executive AI order, and several legislative proposals; and
- Shareholders seek greater transparency into guardrails and their effectiveness in preventing the risks of misinformation and disinformation from generative AI, and the Company's 2024 Responsible AI Transparency Report does not address many of

Board's Perspective

- The Company has a multi-faceted program to address the risks of misinformation and disinformation, including those related to AI;
- Technology companies have a responsibility to help protect democratic processes and institutions globally from cyber-enabled threats, and the Company's Our Democracy Forward Initiative works to safeguard open and secure democratic processes, promote a healthy information ecosystem, and advance corporate civic responsibility;
- In July 2022, the Company completed the acquisition of Miburo, a cyber threat analysis and research company specializing in the detection of and response to foreign information influence, and since then, the Company has created the Microsoft Threat Analysis Center, which brings together analysts focused on nation-state driven cyber and influence actors to identify, analyze, and expose cyber-enabled influence operations targeting democracies around the world;
- The Company recognizes the power of generative AI and the potential that it can be misused, and the Company has implemented a layered approach to address information integrity risks, including developing its foundational commitment to Responsible AI, as well as its Information Integrity Principles;
- The Company is a signatory of the Tech Accord to Combat Deceptive AI in the 2024 Elections and has agreed to eight core commitments to mitigate the impact of deceptive AI on elections around the world; and

these critical questions.

The proponent and Open MIC have filed an [exempt solicitation](#) urging support for this proposal.

- The Company provides several public reports on its efforts to address misinformation and disinformation, including the Company's inaugural Responsible AI Transparency Report, as well as the Microsoft Digital Defense Report and product-specific public reports that describe how it identifies potential risks, measure their propensity to occur, and build mitigations to address.

THE PROPONENT

Arjuna Capital

[Arjuna Capital](#) states that it is a boutique wealth management firm that invests its clients' portfolios for return and impact. It states that, as an "investment manager focused on sustainability [it] understand[s] that social justice, environmental responsibility, and economic vitality are all 'bottom line' issues." and that it is "convinced that investing in a more equitable, environmentally responsible economy is simply smarter long-term investing to secure [clients] wealth and the property of future generations." It states that it avoids and divests from companies whose products and services "are harmful to people and the planet -- think tobacco, weapons or fossil fuels." As of March 2024, it had \$458 million in [AUM](#) and served 155 clients, most of which were high-net-worth individuals.

With regard to [shareholder engagement](#), it states that it filed 16 shareholder proposals in 2022, which focused on a number of issues, including sexual harassment, racist police brutality, board diversity, net zero targets and transition plans, and human and civil rights oversight. It also has a strong focus on issues related to racial and gender pay equity, stating that between 2015 and 2022, it had direct engagement with 35 companies, which had resulted in 32 commitments to disclose pay gaps. In addition, alongside Proxy Impact, Arjuna has developed a [Racial and Gender Pay Scorecard](#) that "provides background on shareholder engagement, regulatory pressure, and the business case for pay equity."

Based on the disclosure provided by companies concerning the identity of proponents, during the first half of 2023, Arjuna submitted eight shareholder proposals to a vote that received 16.4% average support, with none of its proposals receiving majority support.

GLASS LEWIS ANALYSIS

Glass Lewis recommends that shareholders take a close look at proposals such as this to determine whether the actions requested of the Company will clearly lead to the enhancement or protection of shareholder value. Glass Lewis believes that directors who are conscientiously exercising their fiduciary duties will typically have more and better information about the Company and its situation than shareholders. Those directors are also charged with making business decisions and overseeing management. Our default view, therefore, is that the board and management, absent a suspicion of illegal or unethical conduct, will make decisions that are in the best interests of shareholders.

In this case, the proposal is centered on the Company's investment in OpenAI and its ChatGPT chatbot. [OpenAI](#) is an AI research and deployment company that performs research on generative models and offers a range of AI models for consumer and commercial use, though it originally [started](#) as a non-profit artificial intelligence research company that [transitioned](#) to a "capped-profit" company in 2019. The Company states that it has a long-term partnership with OpenAI and that it deploys OpenAI's models across its consumer and enterprise products, adding that it has increased its investments in the development and deployment of specialized supercomputing systems to accelerate OpenAI's research (2024 10-K, p.5). OpenAI's AI models are based on its generative pre-trained transformer ("GPT") technology, which is a type of machine learning model that is pre-trained on a massive amount of data to generate human-like text. GPTs can be further fine-tuned for a range of tasks, such as answering questions, translating language, and summarizing text. (Fawad Ali. " [GPT-1 to GPT-4: Each of OpenAI's GPT Models Explained and Compared](#)." *Make Use Of*. April 11, 2023). Given the nature and scope of the Company's operations, it may be exposed to significant risks arising from its involvement in generative artificial intelligence ("AI"), including, specifically, those related to misinformation and disinformation.

OPENAI AND CHATGPT

Released in 2018, OpenAI's GPT-1 was able to generate fluent and coherent language when given a prompt but was prone to generating repetitive text and was unable to reason over multiple turns of dialogue (Fawad Ali. " [GPT-1 to GPT-4: Each of OpenAI's GPT Models Explained and Compared](#)." *Make Use Of*. April 11, 2023). The Company began [investing](#) in OpenAI in 2019, contributing \$1 billion to support building artificial general intelligence ("AGI") with widely distributed economic benefits. At the time, OpenAI stated that it was partnering with the Company to develop a hardware and software platform within Microsoft Azure which would scale to AGI. Further, it explained that they would jointly develop new Azure AI supercomputing technologies, and the Company would become its exclusive cloud provider, working together to further extend the Company's Azure capabilities in large-scale AI systems. By 2020, the Company was exclusively licensing OpenAI's GPT-3 language model, and it began promoting the possibility for natural language

coding using GPT-3 in 2021 (Mark Labbe. [Microsoft Exclusively Licenses OpenAI's GPT-3 Language Model.](#) *TechTarget*. September 23, 2020; Frederic Lardinois. ["Microsoft Uses GPT-3 to Let Your Code in Natural Language."](#) *TechCrunch*. May 25, 2021). Immediately thereafter, OpenAI and the Company announced a \$100 million startup fund to invest in a relatively small number of early-stage AI companies that were either tackling serious issues that would "benefit all of humanity" or that were addressing productivity improvements (Devin Coldewey. ["OpenAI's \\$100M Startup Fund Will Make 'Big Early Bets' with Microsoft as Partner."](#) *TechCrunch*. May 26, 2021).

Following further GPT development and an additional \$2 billion investment in OpenAI in 2021, OpenAI [announced](#) the launch of its ChatGPT model in November 2022, explaining that its new model could interact "in a conversational way" and that the "dialogue format" made it possible for ChatGPT to "answer followup questions, admit its mistakes, challenge incorrect premises, and reject inappropriate requests." Just two months after launching, the OpenAI model had already acquired an estimated 100 million monthly users, making it the fastest-growing consumer application in history, according to a UBS study (Cade Metz, Karen Weise. ["Microsoft to Invest \\$10 Billion in OpenAI, the Creator of ChatGPT."](#) *The New York Times*. January 23, 2023; Krystal Hu. ["ChatGPT Sets Record for Fastest-Growing User Base - Analyst Note."](#) *Reuters*. February 2, 2023). At the same time, the Company announced in January 2023 that it would make a multiyear, multibillion-dollar investment in OpenAI (Cade Metz, Karen Weise. ["Microsoft to Invest \\$10 Billion in OpenAI, the Creator of ChatGPT."](#) *The New York Times*. January 23, 2023). The following month, the Company [announced](#) that it was launching its new, AI-powered Bing search engine and Edge browser, which would deliver better search, more complete answers, a new chat experience, and the ability to generate content. It further described the new tools as "an AI copilot for the web." The Company also explained that its new Edge browser features could allow a user to ask for summaries of lengthy reports or comparisons and graphics of competing information, among other things. Regarding these new innovations, the Company attributed them to its own commitment to building Azure into an AI supercomputer for the world and to OpenAI's using this infrastructure to train the breakthrough models now being optimized for Bing. Then in March 2023, OpenAI released [GPT-4](#), which is capable of processing both text and images as input and outputting responses with human-like performance on various professional and academic benchmarks. OpenAI reported that its GPT-4 model is able to pass a simulated bar exam with a score around the top 10% of test takers, with similarly high scores in other tests such as the SAT, GRE, and Sommelier certification exams.

In September 2023, OpenAI folded its DALL-E image generator into ChatGPT and released a version of the chatbot that can interact through spoken words, and in November 2023, OpenAI announced that it was offering a service to allow individuals and small businesses to customize their ChatGPT chatbot and instantly share it online, all as part of its ChatGPT Plus monthly subscription. With this customizable option, the Company's vice president of consumer and enterprise product, Peter Deng, stated that because the technology could be used to produce offensive, untruthful, or dangerous material, it would vet all new bots created and disallow those in violation of its terms of service. He also added that users could request that their documents and data not be used to train future versions of their technology, which is already facing lawsuits from writers, artists, and computer programmers for illegal use of their work for AI development (Cade Metz. ["OpenAI Lets Mom-and-Pop Shops Customize ChatGPT."](#) *The New York Times*. November 6, 2023). At the same DevDay showcase event, the Company presented its new version of the chatbot, GPT-4 Turbo, which it described as more capable and able to retrieve information about world and cultural events as recent as April 2023. Further, it discussed a range of other new tools, including its GPT-4V model with a vision tool that can analyze and describe complex images (Barbara ORtutay and Matt O'Brien. ["ChatGPT-make OpenAI Hosts First Big Tech Showcase As It Faces Growing Competition."](#) *ABC News*. November 6, 2023). Continuing the series of rapid developments, the Company also announced that its 365 Copilot AI tool, which integrates the large language model technology from OpenAI's ChatGPT into Office applications, would launch with early access in November 2023. With the announcement, the Company stressed that AI-powered chat data is protected and will not leak outside an organization (Imad Khan. ["Microsoft 365 Copilot AI Tool Will Cost \\$30 Per Month, Launching Nov. 1."](#) *CNET*. October 31, 2023).

OpenAI's CFO explained that, as of late October 2024, 75% of the firm's business was generated from consumer subscriptions, adding that OpenAI has been "wowed" by the pace of growth, especially on the consumer side. In September 2024, OpenAI reached 1 million paid users for its corporate ChatGPT, 250 million weekly active users, and was converting free users to subscription users at a rate of 5-6%. Additionally, OpenAI recently closed a \$6.6 billion fundraising round and asked global banks for a \$4 billion revolving line of credit, to address significant expenses from developing and operating advanced AI systems (Shirin Ghaffary, Edward Ludlow. ["OpenAI CFO Says 75% of Its Revenue Comes from Paying Consumers."](#) *Bloomberg*. October 28, 2024).

OpenAI also announced in late October 2024 that it was working with Broadcom and TSMC to build its own chip to support AI systems, in addition to using AMD chips and Nvidia chips to meet growing demand. Though OpenAI was no longer pursuing plans for its own foundry to build chips, it was focused on in-house chip design efforts. Further, it was still considering whether it would develop or acquire other elements related to its chip design, and may engage additional partners. Meanwhile, through Broadcom, OpenAI secured manufacturing capacity with Taiwan Semiconductor Manufacturing Company to make its first custom-designed chip in 2026, though the timeline could change. Due to the high costs of training AI models and operating services, OpenAI projected a \$5 billion loss in 2024 on \$3.7 billion in

revenue, hence the firm's efforts to diversify suppliers and optimize utilization (Krystal Hu, Fanny Potkin, Stephen Nellis. ["Exclusive: OpenAI Builds First Chip with Broadcom and TSMC. Scales Back Foundry Ambition."](#) *Reuters*. October 30, 2024).

Copilot

The Company has integrated ChatGPT into its AI-powered digital assistant Copilot. In March 2023, the Company [introduced](#) its Microsoft 365 Copilot for work, combining large language models "LLM" with users' data in the Microsoft Graph and Microsoft 365 apps. The Company explained that Copilot would be embedded in Word, Excel, PowerPoint, Outlook, Teams, in addition to providing Business Chat, which would work across the LLM, the Microsoft 365 apps, and users' data such as calendars, emails, chats, documents, meetings, and contacts.

CONCERNS REGARDING THE COMPANY'S AI PRODUCTS AND SERVICES

The proponent of this proposal notes concern regarding risks resulting from the misinformation and disinformation stemming from the Company's strategic partnership with OpenAI and its ChatGPT chatbot. As early as April 2023, various news sources began reporting that the ChatGPT chatbot was generating false information under the guise of their publications. For instance, after hearing from readers trying to track down sources, *The Guardian* announced that it had discovered the chatbot was producing articles under its banner, using real journalists' names and subjects they had a record of covering (Chris Moran. ["ChatGPT Is Making Up Fake Guardian Articles. Here's How We're Responding."](#) *The Guardian*. April 6, 2023).

At the same time, the *Washington Post* reported that ChatGPT was also producing articles under its banner, including one that accused a law professor of sexual harassment on a class trip, even though no trip had occurred and no accusations had been filed against the professor. Additionally, a regional mayor in Australia also threatened to file a defamation suit against ChatGPT for falsely claiming that he had gone to prison on bribery charges. While the senior communications director for the Company stressed that the Company had "developed a safety system including content filtering, operational monitoring and abuse detection to provide a safe search experience" for its users, she also noted that "users are also provided with explicit notice that they are interacting with an AI system." Further complicating the legal outlook, while online services historically have been shielded from liability for content created by third parties, due to Section 230 of the 1996 Communications Decency Act, some experts specializing in intellectual property law imagine that the statute may not prove sufficient if the falsehoods created by AI do "enough damage to warrant a lawsuit" (Pranshu Verma, Will Oremus. ["ChatGPT Invented a Sexual Harassment Scandal and Named a Real Law Prof as the Accused."](#) *The Washington Post*. April 5, 2023).

The proponent also notes that Copilot produced responses that users referred to as "bizarre, disturbing, and in some cases, harmful" (2024 DEF 14A, p.84). In February 2024, the Company investigated social media claims that Copilot produced potentially harmful responses as it appeared to taunt users who suggested they were considering suicide. The Company stated that its investigation found that some conversations were created through "prompt injecting," a technique that permits users to override a LLM, causing it to perform unintended consequences. However, the Company affirmed that it had strengthened its safety filters to help its system detect and block these types of prompts (James Powel. ["Microsoft Investigates Claims of Chatbot Copilot Producing Harmful Responses."](#) *USA Today*. February 28, 2024).

Meanwhile, a Company engineer, Shane Jones, warned in March 2024 that Copilot Designer can create violent, sexual images and ignore copyrights. The engineer, who worked at the Company for six years, had been testing the AI image generator in his free time and told *CNBC* that he was disturbed by his findings. Jones also stated that the Company wasn't taking appropriate action, and so he escalated the matter by sending letters to FTC chair Lina Khan and to the Company's board. In his letter to the board, Jones requested that the Company's environmental, social, and public policy committee investigate certain decisions by the legal department and management, as well as begin an "an independent review of Microsoft's responsible AI incident reporting processes." Jones further told the board that he's "taken extraordinary efforts to try to raise this issue internally" by reporting concerning images to the Office of Responsible AI, publishing an internal post on the matter, and meeting directly with senior management responsible for Copilot Designer. While the Company acknowledged his concerns, it was unwilling to take the product off the market (Hayden Field. ["Microsoft Engineer Warns Company's AI Tool Creates Violent, Sexual Images, Ignores Copyrights."](#) *CNBC*. March 7, 2024).

Shortly after public reports on these matters, the U.S. House of Representatives set a strict ban on congressional staffers' use of Copilot after it was deemed by the Office of Cybersecurity to be a risk to users due to the threat of leaking House data to non-House approved cloud services. In response, the Company announced a roadmap of AI tools, like Copilot, that meet federal government security and compliance requirements which the Company intended to deliver later in 2024 (Mrinmay Dey, Mehnaz Yasmin. ["US Congress Bans Staff Use of Microsoft's AI Copilot. Axios Reports."](#) *Reuters*. March 29, 2024).

Misinformation and Disinformation Risks from GPT Technology

Despite its capabilities, the responses produced by GPT and other generative AI systems are not entirely reliable. Because AI models, such as GPT, are programmed to guess the next word in a sequence of words based on large volumes of text sourced online, such models are unable to evaluate whether a statement is true, and they are known to respond with false or inaccurate information, a phenomenon referred to as "hallucination" (Karen Weise, Cade Metz. "[When A.I. Chatbots Hallucinate.](#)" *The New York Times*. May 9, 2023). While OpenAI and other AI developers state they are working to make their models more truthful, some experts remain doubtful that generative AI's tendency to hallucinate information can be completely removed due to the technology's nature. The risk of AI hallucination can be more severe depending on the task an AI model is used for. For example, users have employed AI models for high-stakes tasks where accuracy is essential, such as psychotherapy, writing legal briefs, and news-writing, which the director of the University of Washington's Computational Linguistics Laboratory describes as a "mismatch between the technology and the proposed use cases" (Matt O'Brien. "[Chatbots Sometimes Make Things Up. Is AI's Hallucination Problem Fixable?](#)." *Associated Press*. August 1, 2023).

Researchers quickly began examining the potential to spread misinformation using ChatGPT soon after it came on the market, voicing strong concerns, including that chatbots could make disinformation cheaper and easier to produce and spread, and they could also make disinformation more difficult to detect by smoothing out human errors such as poor syntax or mistranslations. For example, a 2019 paper from OpenAI researchers cited concerns that AI "capabilities could lower costs of disinformation campaigns" and aid in the malicious pursuit of "monetary gain, a particular political agenda, and/or a desire to create chaos or confusion." When tested by researchers at the Center on Terrorism, Extremism, and Counterterrorism at the Middlebury Institute of International Studies, GPT-3 presented an "impressively deep knowledge of extremist communities" and could be prompted to generate responses mimicking content from mass shooters, forums discussing Nazism, writings in defense of QAnon, and also multilingual extremist material. Over time, however, ChatGPT showed signs of resisting some researchers' efforts to generate false information, sometimes adding disclaimers to long-debunked arguments. Nevertheless, another research group providing cyber threat intelligence, Check Point Research, stated that cybercriminals were already working with ChatGPT to create malware (Tiffany Hsu, Stuart A. Thompson. "[Disinformation Researchers Raise Alarms about A.I. Chatbots.](#)" *The New York Times*. February 8, 2023).

Although OpenAI and other companies engaged in AI have developed safeguards to prevent their models from generating hate speech and disinformation, researchers have discovered methods to bypass those preventative measures. Such bypasses add to concerns that generative AI models could be used to generate nearly unlimited amounts of harmful information (Cade Metz. "[Researchers Poke Holes in Safety Controls of ChatGPT and Other Chatbots.](#)" *The New York Times*. July 27, 2023). Research from Princeton, Virginia Tech, Stanford, and IBM also found that by using a paid service offered by OpenAI to adjust its GPTs for particular tasks, users were able to bypass guardrails and generate material including political messages, hate speech, and language involving child abuse (Cade Metz. "[Researchers Say Guardrails Built Around A.I. Systems Are Not So Sturdy.](#)" *The New York Times*. October 19, 2023).

A recent [study](#) by Stanford RegLab and Institute for Human-Centered AI researchers showed evidence that popular chatbots from OpenAI, Google, and Meta were prone to hallucinations when answering legal questions, which would pose problems for those using the technology if they could not afford human lawyers. The study reported that when answering questions regarding a court's core ruling, the chatbots would hallucinate at least 75% of the time, based on testing with more than 200,000 legal questions posed to general-purpose AI models. However, the models were more accurate when asked questions regarding cases from the U.S. Supreme Court and the U.S. Courts of Appeal for the Second Circuit and the Ninth Circuit, as those cases may have been more frequently cited and discussed and thus appeared more frequently in the models' training data. The study also revealed that the chatbots were likely to believe a users' false premise embedded in their questions, and subsequently provided answers that reinforced the users' mistakes (Isabel Gottlieb, Isaiah Poritz. "[Popular AI Chatbots Found to Give Error-Ridden Legal Answers.](#)" *Bloomberg Law*. January 12, 2024).

Hallucinations have persisted as a problem for chatbots. In May 2024, Google's chatbot told some users that it was safe to eat rocks, leading the firm to review its AI overviews search feature. In June 2023, two lawyers received fines of \$5,000 after one of them admitted that he had used ChatGPT to help write a court filing, but the chatbot added fake citations for cases that never existed. In response to these ongoing issues, a recent study by the scientific journal *Nature* examined a new method for determining when an AI tool is likely to be hallucinating, demonstrating its ability to accurately discern correct and incorrect AI-generated answers approximately 79% of the time. Currently, the method only addresses one of several causes of AI hallucinations and requires nearly ten times the computing power of a standard chatbot conversation, but researchers were hopeful that it could lead to more reliable AI systems in the near future (Billy Perrigo. "[Scientists Develop New Algorithm to Spot AI 'Hallucinations'.](#)" *Time*. June 19, 2024).

There have also been a number of attempts to curb election-related misinformation and disinformation. For example, 20 tech companies, including the Company, OpenAI, Adobe, Amazon, Anthropic, Meta, Google, TikTok, and X, signed a voluntary [pledge](#) at the Munich Security Conference to prevent deceptive AI content from disrupting voting in 2024. As part of their pledge, the companies agreed to collaborate on AI detection tools but did not ban election-related AI content (Tiffany Hsu, Cade Metz. "[In Big Election Year, A.I.'s Architects Move Against Its Misuse.](#)" *The New York Times*. February

16, 2024). In addition, Google announced that it would restrict its Gemini AI chatbot from responding to questions about global elections in the U.S., South Africa, and India in 2024 in an attempt to avoid problems regarding and fake news from its generative AI. The chatbot was trained to respond that it was still learning how to answer such a question and that users should try Google Search to find information about elections (Zaheer Kachwala. "[Google Restricts AI Chatbot Gemini From Answering Queries on Global Elections](#)." *Reuters*. March 12, 2024).

LAWSUITS INVOLVING THE COMPANY, OPENAI, AND CHATGPT

Section 230 of the Communications Decency Act

As part of its rationale for this proposal, the proponent states that legal experts believe technology companies' liability shield provided under Section 230 of the Communications Decency Act ("CDA 230"), which has historically provided legal protection when third-party content is posted, may not apply to content generated by AI.

Although not directly related to the Company, OpenAI, or ChatGPT, but nevertheless still relevant to this issue, is the March 2024 court ruling *Patterson et al. v. Meta Platforms, Inc. et al.* In March 2024, a New York State trial court ruled in *Patterson* that the CDA 230 does not preclude claims asserted against various social media companies, including Alphabet, Meta, Snap, Discord, Reddit, and Amazon, relating to a racially motivated 2022 mass shooting at a supermarket in Buffalo, New York. In May 2023, a victim, estate representatives, and heirs of victims of a mass shooting at Topps Friendly Market in Buffalo sued various social media companies alleging that the "defective and unreasonably dangerous design" of the social media companies' "defective products", which the lawsuit identified as including Instagram, Facebook, YouTube, Twitch, Snapchat, Discord, and Reddit, facilitated the shooter's conduct. The plaintiffs asserted numerous causes of action against the social media companies, including strict product liability, negligence, negligent failure to warn, unjust enrichment, and infliction of emotional distress. However, the companies attempted to dismiss the claims by asserting, among other things, that CDA 230 immunized them from liability. It appears to be the first time a New York state court determined that the CDA 230 does not prevent claims against social media companies resulting from a mass shooting. Nevertheless, the social media companies filed notices of appeal (Samantha Katze. "[Communications Decency Act Does Not Bar Claims Against Social Media Companies Due to Mass Shooting](#)." *JDSupra*. March 27, 2024).

In a recent article on liability and AI output, UCLA law professor and First Amendment scholar Eugene Volokh shared conclusions following a virtual symposium on AI and free speech, presenting the consensus that Section 230 of the CDA of 1996 does not apply to AI. Volokh suggested that two possible frameworks could enable successful litigation against an AI company: (i) proof of "reckless disregard for the truth" articulated in the 1964 U.S. Supreme Court Case *New York Times v. Sullivan* (in which case, an AI company could be found guilty of "reckless disregard for the truth" if it were alerted that their program was generating specific false and libelous content and took no action); and (ii) proof that flaws in the product design caused it to generate defamatory content. However, the theories had not yet been tested in court (Joel Simon. "[Can AI Be Sued for Defamation?](#)" *Columbia Journalism Review*. March 18, 2024).

Defamation Lawsuits Against Open AI and ChatGPT

Mark Walters, a radio host from Georgia, sued OpenAI in June 2023, claiming that its chatbot generated a legal complaint accusing him of embezzling money from a gun rights group, allegations Walters said he had never faced. The complaint was generated in response to a journalist's research on a real court case but provided information that was entirely fake. Nevertheless, legal experts anticipated a defense could draw on CDA 230 (Isaiah Poritz. "[First ChatGPT Defamation Lawsuit to Test AI's Legal Liability](#)." *Bloomberg Law*. June 12, 2023). The journalist whose queries generated the false information in a ChatGPT hallucination never actually published the misinformation, nor did Walters notify OpenAI that ChatGPT was making false statements about him, two factors which complicate the suit's chances of success (Debra Cassens Weiss. "[Radio Host Faces Hurdles in ChatGPT Defamation Suit](#)." *ABA Journal*. June 12, 2023). In January 2024, a judge in Gwinett County Superior Court in Georgia denied OpenAI's request to dismiss the suit, though no explanation was provided (Isaiah Poritz. "[OpenAI Fails to Escape First Defamation Suit from Radio Host](#)." *Bloomberg Law*. January 16, 2024).

Privacy Lawsuits Against OpenAI and ChatGPT

In September 2023, two unnamed software engineers filed a class-action suit against the Company and OpenAI, claiming that they were training their AI chatbots using hundreds of millions of internet users' stolen personal information from social media platforms and other sites (Blake Brittain. "[OpenAI, Microsoft Hit with New US Consumer Privacy Class Action](#)." *Reuters*. September 6, 2023). A separate suit argues that OpenAI is not transparent enough with its users so they understand that the data they put into the model may be used to train new products from which it will generate revenues. It also claims that OpenAI does not do enough to ensure children under 13 are not using its tools (Gerrit De Vynck. "[ChatGPT Maker OpenAI Faces a Lawsuit over How It Used People's Data](#)." *The Washington Post*. June 28, 2023). The suit also [refers](#) to the potential for AI hallucinations as a cause for concern regarding the chatbots' use of users' private data (p.50).

Copyright Infringement Lawsuits Against OpenAI and ChatGPT

Multiple authors have filed at least three lawsuits against OpenAI for using their copyrighted work to train the ChatGPT chatbot, including a lawsuit that lists the Authors Guild and famous writers such as Johnathan Franzen, John Grisham, George RR Martin, and Jodi Picoult, among others, as plaintiffs (Blake Montgomery. " [OpenAI Offers to Pay for ChatGPT Customers' Copyright Lawsuits.](#)" *The Guardian*. November 6, 2023). The main argument of the suit is that "the ingestion of works of authorship as training data is itself a reproduction of the works." Moreover, the case argues that when OpenAI scans and uses writers' content without copyright permission, it helps foster work that publishers would otherwise have to pay authors to create. What makes the suit so complicated is the difficulty in determining what the exact data sets used by OpenAI to train its chatbots and how they work (Max Zahn. " [Authors' Lawsuit Against OpenAi Could 'Fundamentally Reshape' Artificial Intelligence. According to Experts.](#)" *ABC News*. September 25, 2023).

Some legal experts speculated that this suit could fundamentally impact the direction and capabilities of generative AI, though it has yet to be determined whether the outcome will create new limitations on the core processes of the technology or whether it will enable the current, expansive approach to online material to continue. OpenAI and other companies facing litigation for their generative AI, such as Meta, are expected to rely on more open-ended concepts of "fair use" in their defense. In fact, in response to a similar suit brought by comedian Sarah Silverman, lawyers for Meta stated in their defense filing that "copyright law does not protect facts or the syntactical, structural, and linguistic information that may have been extracted from books like the Plaintiffs' during training." Further, they added that the "use of texts to train" chatbots to "statistically model language and generate original expression is transformative by nature and quintessential fair use." Attorneys for OpenAI stated in their filings that the plaintiffs' copyright claims "misconceive the scope of copyright, failing to take into account the limitations and exceptions (including fair use) that properly leave room for innovations like the large language models now at the forefront of artificial intelligence" (Max Zahn. " [Authors' Lawsuit Against OpenAi Could 'Fundamentally Reshape' Artificial Intelligence. According to Experts.](#)" *ABC News*. September 25, 2023).

In February 2024, a California court ruled to dismiss five of the six claims made in the suit brought by Silverman and authors Christopher Golden, Richard Kadrey, Paul Tremblay, and Mona Awad. The judge did not agree with the plaintiffs' claims that OpenAI had committed unlawful business practices and fraudulent conduct related to unfair competition, though it upheld the unfair competition claim that the firm did not request permission to use their work for commercial profit. Additionally, the judge stated that the authors' claim regarding "risk of future damage to intellectual property" was too speculative to consider. OpenAI had asked the court to dismiss all counts but the first and main complaint: direct copyright infringement. The firm was also facing several copyright infringement lawsuits from other authors, including a proposed class action lawsuit (Emilia David. " [Sarah Silverman's Lawsuit Against OpenAI Partially Dismissed.](#)" *The Verge*. February 13, 2024).

The New York Times also sued OpenAI and the Company in late December 2023 for copyright infringement. While the lawsuit does not include an exact monetary request, it claims the defendants should be held responsible for "billions of dollars in statutory and actual damages" related to the "unlawful copying and use of The Times's uniquely valuable works." It also asks that both companies destroy any chatbot models and training data that use copyrighted material from *The New York Times* (Michael M. Grynbaum, Ryan Mac. " [The Times Sues OpenAI and Microsoft Over A.I. Use of Copyrighted Work.](#)" *The New York Times*. December 27, 2023). Two months after the lawsuit was filed, OpenAI requested that some parts of the suit be dismissed, claiming that *The New York Times* had "hacked" its chatbot and other AI systems to generate misleading evidence for the case. The firm argued that the paper used "deceptive prompts" that violated OpenAI's terms of use and that it had paid someone to hack OpenAI's products. However, OpenAI did not name the hacker or identify any anti-hacking laws that the paper allegedly broke. A lawyer for the *The New York Times* responded that the paper had not hacked anything but had simply looked for evidence that the chatbot had stolen and reproduced the paper's copyrighted content (Blake Brittain. " [OpenAI Says New York Times 'Hacked' ChatGPT to Build Copyright Lawsuit.](#)" *Reuters*. February 27, 2024).

Further, a group of eight additional U.S. newspapers (*New York Daily News*, *Chicago Tribune*, *Denver Post*, *Mercury News*, *Orange Country Register*, *St. Paul Pioneer-Press*, *Orlando Sentinel* and *South Florida Sun Sentinel*) filed a lawsuit against the Company and OpenAI, also claiming that the companies stole millions of copyrighted news articles to train their chatbots, without permission or payment. OpenAI stated that it was unaware of the papers' owner Alden Global Capital's concerns, but added that it was "actively engaged in constructive partnerships and conversations with many news organizations around the world to explore opportunities, discuss any concerns, and provide solutions." Although tech companies have argued that the "fair use" doctrine enables them to use publicly accessible internet content to train their chatbots, some companies have established deals to pay for content to avoid potential legal challenges. For example, OpenAI entered into a partnership with the *Associated Press* ("AP") in 2023 to pay an undisclosed amount to license the AP's archive of news stories. OpenAI also made licensing deals with publishing firms *Axel Springer* in Germany, *Prisa Media* in Spain, *Le Monde* in France, and the *Financial Times* in the UK. (" [Eight US Newspapers Sue ChatGPT-Maker OpenAI and Microsoft for Copyright Infringement.](#)" *AP News*. April 30, 2024).

FEDERAL TRADE COMMISSION SCRUTINY

AI systems have come under regulatory scrutiny as well. In July 2023, the Federal Trade Commission ("FTC") opened an investigation into whether OpenAI's ChatGPT tool harmed consumers through its collection of data and publication of false information on individuals. As part of the investigation, the FTC asked OpenAI about its AI model training and its use of personal data, and demanded that the firm provide it with documents and details (Cecilia Kang, Cade Metz. "[F.T.C. Opens Investigation Into ChatGPT Maker Over Technology's Potential Harms](#)." *The New York Times*. July 13, 2023).

In January 2024, the FTC began an inquiry into multibillion-dollar investments by the Company, Amazon, and Google in OpenAI and Anthropic. (The Company committed \$13 billion for a 49% stake in OpenAI; Amazon said it would invest up to \$4 billion in Anthropic, while Google committed to investing more than \$2 billion in Anthropic.) The FTC would be examining how such investment deals altered competition in the industry, potentially informing federal antitrust investigations regarding whether deals had broken the law. The FTC planned to ask the Company, OpenAI, Amazon, Google, and Anthropic to describe their influence over their partners as well as how they worked together when making decisions. It would also be asking whether any of the deals between the tech companies and the startups have rights to board seats or other oversight arrangements over each other. While the Company obtained a seat on OpenAI's board in November 2023, it did not have the right to vote on its decisions. According to FTC chair Lina Khan, the FTC study would "shed light on whether investments and partnerships pursued by dominant companies risk distorting innovation and undermining fair competition." In recent years, the FTC has divided investigations of corporate mergers with the Justice Department, such that the FTC filed antitrust lawsuits against Amazon and Meta while the Justice Department sued Google and was investigating Apple (David McCabe. "[Federal Trade Commission Launches Inquiry into A.I. Deals by Tech Giants](#)." *The New York Times*. January 25, 2024).

Similarly, the Competition and Markets Authority, a British regulator, announced in December 2023 that it was investigating the Company's deal with OpenAI, and the European Commission was considering whether antitrust laws could apply. The Company's competition lawyer stated that the U.S. had "assumed a global A.I. leadership position because important American companies are working together," and that the Company's partnership with OpenAI was "promoting competition and accelerating innovation" (David McCabe. "[Federal Trade Commission Launches Inquiry into A.I. Deals by Tech Giants](#)." *The New York Times*. January 25, 2024).

CONCERNS REGARDING THE SPREAD OF MISINFORMATION AND DISINFORMATION ONLINE

One of the most common criticisms of social media in recent years has been the dissemination of misinformation and disinformation. This issue first attracted global scrutiny in 2016, as the deliberate dissemination of false information was employed as a tactic to defeat political rivals during the U.S. Presidential election. In 2017, the non-governmental organization Freedom House released a [report](#) that found that "online manipulation and disinformation tactics played an important role" in elections in 18 countries, including the U.S., and noted Russia's influence on the U.S. election. According to this report, internet freedom declined for seven consecutive years, largely as a result of disinformation tactics. Specifically, the report stated that "'fake news,' automated 'bot' accounts, and other manipulation methods gained particular attention in the United States," and that, although the online environment in the U.S. "remained generally free, it was troubled by a proliferation of fabricated news articles, divisive partisan vitriol, and aggressive harassment of many journalists, both during and after the presidential election campaign." The report noted that around 30,000 fake accounts were removed from Meta's platform before the 2017 French elections. The report further noted how preventing content manipulation could be difficult given its decentralized nature and the rampant employment of people and bots for such purposes.

Further, in March 2018, *Science* published a study that analyzed every major contested news story in English since the Twitter platform came into existence, comprising 126,000 stories tweeted by 3 million users over more than 10 years. The study [found](#) that "[f]alsehood diffused significantly farther, faster, deeper, and more broadly than the truth in all categories of information[.]" However, it was also found that bots amplified true stories as much as false ones, leading the researchers to conclude that people are ultimately still responsible for the propagation of fake news. Ultimately, it demonstrated that social media serves to systematically amplify falsehood over truth (Robinson Meyer. "[The Grim Conclusions of the Largest-Ever Study of Fake News](#)." *The Atlantic*. March 8, 2018). One of the study's researchers recognized the challenge that such polarization poses to garnering support from platforms such as that operated by Meta, given that polarization "has turned out to be a great business model" (Steve Lohr. "[It's True: False News Spreads Faster and Wider. And Humans Are to Blame](#)." *New York Times*. March 8, 2018).

With the emergence of AI technology, misinformation can now also be spread through chatbots that can generate false information. When Alphabet introduced its Bard chatbot with a live presentation on Twitter in February 2023, the chatbot provided inaccurate information about which satellite first took photographs outside the Earth's solar system. Bard responded that the James Webb Space Telescope took the first pictures, when it was actually the European Southern Observatory's Very Large Telescope. Alphabet responded to the error by stating that it "highlights the importance of a rigorous testing process, something that [it] is kicking off this week." However, the firm had lost \$100 billion in market value, with Alphabet's shares declining as much as 9% following the episode and concerns about competition from the

Company. At the time, the Bard chatbot ad had been viewed more than one million times (Martin Coulter, Greg Bensinger. "[Alphabet Shares Dive After Google AI Chatbot Bard Flubs Answer in Ad.](#)" *Reuters*. February 9, 2023).

Chabots can also generate deepfake photographs, videos, and audio recordings. Much of the attention devoted to the harms of deepfake material relates to pornographic material and "sextortion" cases, including materials involving children. Researchers noted in 2019 that 96% of the 14,000 deepfake videos online were pornographic, with other reports indicating that as many as 143,000 videos appearing on 40 of the most popular websites for faked videos had been viewed more than 4.2 billion times. The rapid increase in the creation of deepfakes, and especially those involving minors, has prompted responses from attorneys general in all 50 states, who asked Congress to strengthen their tools for fighting AI child sexual abuse images. Legislators around the country are also beginning to develop legislation to address the harms and repercussions of such AI-generated materials. This legislation has taken a wide range of approaches, including allowing victims to pursue civil lawsuits, criminalizing some conduct, and considering steps such as mandating digital watermarks to help trace images and content back to their producers (Frank Figliuzzi. " [A Loophole Makes It Hard to Punish These Despicable AI-Generated Nude Photos.](#)" *MSNBC*. November 7, 2023).

Election Interference

The issue of misinformation and disinformation generated by AI has become increasingly salient with respect to election interference. For example, despite calls for caution, a recent report found that OpenAI had not been enforcing its rules regarding political campaigns' use of its AI tools. When OpenAI first launched ChatGPT in November of 2022, it banned political campaigns from using the chatbot as a safeguard against potential election risks. Further, in March 2023, when OpenAI updated the rules on its website, it banned political campaigns specifically from using ChatGPT to generate content targeted to specific demographics, a practice that could spread disinformation at an unprecedented scale. Despite this policy, however, it was reported in August 2023 that OpenAI hadn't enforced its ban for months. Many researchers and politicians have raised alarms about the potential for election interference and the ability for fake photos and videos to mislead voters. In response to the need to block the most worrying ways the chatbot could be used in politics, OpenAI used its ChatGPT to create novel political risks for review, and then it developed a policy that prohibits "scaled uses" for political campaigns or lobbying. In response to this issue, a representative for OpenAI's product policy stated that it was "building out... safety capabilities" and that it was exploring ways to detect generated content related to campaign materials. However, OpenAI admits that the "nuanced" nature of its rules, which it developed to accommodate the benefits and address the risks of political content, can make enforcement of its policies difficult (Cat Zakrzewski. " [ChatGPT Breaks Its Own Rules on Political Messages.](#)" *The Washington Post*. August 28, 2023).

Concerns about the impact of AI-generated deepfakes on politics have been growing, as more fake videos featuring major political figures have begun to circulate on social media. For example, during the U.S. presidential primary election in June 2023, Ron DeSantis's campaign began using three images of Donald Trump embracing Dr. Anthony Fauci in a campaign video, which the agency Agence France-Presse first reported to be deepfake images. DeSantis did not respond, but Trump posted on social media criticizing the use of the AI images as unacceptable, after which the DeSantis campaign responded that they were obviously fake and intended to be humorous (Nicholas Nehamas. "[DeSantis Campaign Uses Apparently Fake Images to Attack Trump on Twitter.](#)" *The New York Times*. June 8, 2023).

Many political lawyers have expressed concerns that the lack of federal regulation regarding AI as well as insufficient campaign funds could leave many politicians and political groups unable to defend against targeted fake content. One factor complicating legal strategies to counter deepfakes is that such AI-generated false content is often very difficult to trace. In May 2023, the American Association of Political Consultants, a bipartisan board, unanimously agreed to condemn campaigns' use of generative deepfakes, and they also encouraged the media to refuse to carry or deliver ads using deepfakes. Additionally, the Federal Election Commission solicited public feedback in August 2023 on possible regulations to address deepfakes in political ads. There have also been a number of other regulatory and legislative responses to this issue. For example, Representative Yvette Clark of New York introduced a bill to require disclosure in AI-generated political ads, and several states have passed AI regulations. Nevertheless, legal experts say that necessary regulation will not likely catch up until a particularly bad deepfake creates a strong reaction, and it won't come before the 2024 election. Additionally, experts expressed concern regarding the impact of AI deepfakes on the ability of people to know the truth and on democracy, and they emphasized that foreign political actors could also utilize deepfakes against the U.S. (Tatyana Monnay. " [Deepfake Political Ads Are 'Wild West' for Campaign Lawyers.](#)" *Bloomberg Law*. September 5, 2023).

After Google stated that it would impose new labels on deceptive AI-generated political advertisements that could fake a candidate's voice or actions, U.S. regulators called on social media platforms Facebook, Instagram, and X to explain why they have not implemented similar measures. In October 2023, two members of Congress sent a letter to the CEOs of Meta and X expressing "serious concerns" about the emergence of AI-generated political ads on their platforms and asking each to explain any rules they're creating to curb the harms to free and fair elections. The letter to the executives warns that, particularly in light of the 2024 elections, "a lack of transparency about this type of content in political ads could lead to a dangerous deluge of election-related misinformation and disinformation across your platforms – where

voters often turn to learn about candidates and issues" (Matt O'Brien. "[Meta and X Questioned By Lawmakers Over Lack of Rules Against AI-generated Political Deepfakes](#)." *AP News*. October 5, 2023).

These concerns are even more salient given actual instances of the use of AI in electioneering activities. For example, in January 2024, the New Hampshire attorney general's office announced that it would be investigating reports of a robocall using an AI-generated voice to mimic President Joe Biden that discouraged people from voting in the state's primary election. The attorney general stated that the calls appeared to be an attempt to suppress voting, incorrectly telling voters that they needed to save their vote for the November election (Ali Swenson and Will Weissert. "[New Hampshire Investigating Fake Biden Robocall Meant to Discourage Voters Ahead of Primary](#)." *AP News*. January 22, 2024).

Researchers and bipartisan election officials also tested the accuracy of information provided by chatbots regarding the upcoming U.S. primary elections at Columbia University in February 2024. Their results showed that both GPT-4 and the Google's Gemini tended to provide information that was illogical and included outdated responses, and in some cases, would even send voters to nonexistent polling locations. All five models (OpenAI's GPT-4, Google's Gemini, Meta's LLaMA 2, Anthropic's Claude, and Mixtral from the French company Mistral) tested by the group were unable to provide answers to basic questions about the election process, with participants rating more than half of the chatbots' responses as inaccurate and 40% categorized as harmful, including perpetuating dated and inaccurate information that could limit voting rights. For example, when asked about a majority Black neighborhood in Philadelphia, Gemini responded that there was no voting precinct in the U.S. for that zip code. Overall, Google's Gemini got almost two-thirds of all answers wrong, while LLaMA 2 and Mixtral also had high rates of incorrect information (Garance Burke. "[Chatbots' Inaccurate, Misleading Responses About US Elections Threaten to Keep Voters from Polls](#)." *AP News*. February 27, 2024).

In February 2024, FBI Director Christopher Wray warned that the U.S. expected fast-moving threats, and more of them, as AI and other technological advancements had made election interference much easier than before (Eric Tucker. "[The US Is Bracing for Complex, Fast-Moving Threats to Elections this Year, FBI Director Warns](#)." *AP News*. February 29, 2024). Europe and Asia had already experienced a wave of AI deepfakes for months, with examples of AI deepfakes seen in: (i) a video of Moldova's pro-Western president supporting a Russia-friendly political party; (ii) audio clips of the liberal party leader in Slovakia talking about rigging votes; and (iii) a video of an opposition lawmaker in Bangladesh wearing a bikini in the Muslim majority nation (Ali Swenson, Kelvin Chan. "[Election Disinformation Takes a Big Leap with AI Being Used to Deceive Worldwide](#)." *AP News*. March 14, 2024).

Meanwhile, a study in Europe found that popular AI chatbots weren't providing users with accurate information about upcoming elections. The chatbots, including the Company's Copilot, OpenAI's ChatGPT 3.5 and 4.0, as well as Google's Gemini, were tested by the Berlin-based nonprofit Democracy Reporting International on 400 election-related questions posed in ten different languages and relating to the election process in 10 EU countries, with questions written in simple language fit for the average user of these AI chatbots. The results showed that none of the four chatbots were able to "provide reliably trustworthy answers." Further, when asked about a topic without a lot of information, the chatbots tended to invent information to supply answers. Google's Gemini was found to provide the most misleading or false information, and it had the highest number of refusals to answer questions (Anna Desmarais. "[Chatbots Providing 'Unintentional' Misinformation Ahead of EU Elections](#)." *EuroNews*. April 22, 2024).

Further, a study by data analytics startup GroundTruthAI reported that, based on 2,784 responses, large language models, including ChatGPT and Alphabet's Gemini gave inaccurate information 27% of the time when asked about the 2024 election and voting. According to the study, the five models it examined answered correctly 73% of the time, but OpenAI's latest version, GPT-4o, answered correctly 81% of the time, while Google's Gemini 1.0 Pro answered correctly just 57% of the time. GroundTruthAI's CEO stated that the study should serve as a warning to companies incorporating AI into their search functions. The same mistrust of chatbots as a source of election information could also be seen in a recent survey from Pew Research Center, which found that although more Americans were using ChatGPT, approximately 40% of adults surveyed stated they don't trust the information that comes from the chatbot regarding the 2024 presidential election (Aaron Franco, Morgan Radford. "[AI Chatbots Got Questions About the 2024 Election Wrong 27% of the Time, Study Finds](#)." *NBC News*. June 7, 2024).

The news monitoring service NewsGuard reported in June 2024 that generative AI models, including ChatGPT, Copilot, Meta AI, Google Gemini, You.com's Smart Assistant, Grok, Inflection, Mistral, Anthropic's Claude, and Perplexity, were spreading Russian disinformation. Specifically, NewsGuard's study found that when prompted, the ten chatbots spread Russian disinformation narratives 32% of the time. The study employed 57 prompts based on 19 significant false narratives that the news monitoring service had linked to the Russian disinformation network, including claims about Ukraine President Volodymyr Zelenskyy's corruption. NewsGuard submitted its study to the U.S. AI Safety Institute of the National Institute of Standards and Technology and the European Commission. At the same time, however, the U.S. House Committee on Oversight and Accountability launched an investigation into NewsGuard over its "potential to serve as a non-transparent agent of censorship campaigns," an accusation which NewsGuard rejected (Pascale Davies. "[ChatGPT, Grok, Gemini and other AI Chatbots Are Spewing Russian Misinformation, Study Finds](#)." *Euro News*. June 18, 2024).

Copilot and \ChatGPT also received criticism in late June 2024 when, hours before the first presidential debate, they repeated misinformation that CNN's broadcast of the debate would occur with a "1-2 minute delay." When prompted by NBC News, the chatbots neglected to include CNN's denial of the false claim, with Copilot instead citing the website of former Fox News host Lou Dobbs, which had cited a conservative writer's post about an alleged 1-2 minute delay. ChatGPT cited Katie Couric Media and UPI even though neither of their original articles mentioned that there would be a delay during CNN's broadcast of the debate. Meanwhile, NBC reported that Meta AI and X's Grok answered the questions correctly, but Google's Gemini declined to answer the questions about whether there would be a delay, which it deemed too political (Ben Goggin. "[OpenAI's ChatGPT and Microsoft's Copilot Repeated a False Claim About the Presidential Debate.](#)" NBC News. June 28, 2024).

In late October 2024, the FBI announced that a video appearing to show someone destroying mail-in ballots for Donald Trump in Pennsylvania had been created and amplified by Russian actors. U.S. officials from the FBI, the Office of the Director of National Intelligence, and the Cybersecurity and Infrastructure Security Agency stated that the fake video was part of "Moscow's broader effort to raise unfounded questions about the integrity of the U.S. election and stoke divisions among Americans." A day prior to the announcement from federal officials, the Bucks County Board of Elections in Pennsylvania identified the video as fake, stating that the materials in the video were clearly not authentic materials distributed by the board. The video appeared to show a Black person with a foreign accent tearing up ballots marked for Trump, which some experts believed was an intentional choice to amplify racism with disinformation (Melissa Goldin, Mike Catalini, Ali Swenson. "[Russian Actors Made Fake Video Depicting Mail-In Ballots for Trump Being Destroyed.](#) FBI Says." AP News. October 25, 2024).

The following week, just days before the U.S. presidential election, the FBI announced that "Russian influence actors" were the source of a video depicting alleged voter fraud in Georgia. The video in Georgia showed a person who identified himself as a Haitian immigrant who planned to vote multiple times in two different Georgia counties for Vice President Kamala Harris. This person briefly showed several Georgia IDs with different names and addresses, but neither of the IDs contained information matching any registered voter in the two counties in question, Gwinnett and Fulton. In response, Georgia Secretary of State Brad Raffensperger stated that the video was "obviously fake" and most likely the product of Russian trolls "attempting to sow discord and chaos on the eve of the election." U.S. intelligence officials echoed that finding the following day, pointing to "Russian influence actors" and explaining that the intelligence community expected Russia to continue "to create and release additional media content that seeks to undermine trust in the integrity of the election and divide Americans." While the original post circulating the Georgia video was no longer available online, people were still widely sharing various copycat versions claiming election fraud (Eric Tucker, Ali Swenson. "[FBI Links Video Falsely Depicting Voter Fraud in Georgia to 'Russian Influence Actors'.](#)" AP News. November 1, 2024).

Raffensperger directly urged social media companies to remove the fake Georgia video from their platforms. While a spokesperson for X responded that it was "taking action against the posts," which violated its policies, the video was still visible in posts on X two hours later. Meta stated that it was adding a label to the bottom of the video anytime it appeared on its platform stating that the content was "manufactured by Russian influence actors" (Rami Ayyub, Susan Heavey, Katie Paul, Sheila Dang. "[US Blames Russia over Video Falsely Alleging Fraudulent Voting in State of Georgia.](#)" Reuters. November 1, 2024).

Misinformation Related to Hurricane Helene

Following the devastation of Hurricane Helene in early October 2024, victims were confronted with rapidly spreading misinformation regarding recovery efforts. Many local media outlets worked nonstop to counter false information and AI-generated images that emerged after the disaster, but the amount of rumors and theories was so great, it caused residents to inundate local officials with questions about the misinformation and prevented access to victims in need. Some of the disinformation was being spread by elected officials in the areas affected, as well as by former President Trump, who repeated rumors that the federal government was confiscating and diverting aid meant for the victims of Helene. The rumors prompted North Carolina State Senator Kevin Corbin to post online requesting followers to "help STOP this conspiracy theory junk that is floating all over Facebook and the internet about the floods in [Western North Carolina]," claiming the rumors were "just a distraction from people trying to do their job." Similarly, the American Red Cross posted online that misinformation was interfering with its ability to deliver critical aid to those affected. North Carolina Governor Roy Cooper said at a press conference that the rumors were having a real impact on the recovery efforts, while FEMA's administrator stated that "[t]his level of misinformation creates the scenario where [the victims] won't even come to us.... won't even register.... so they can get what they're eligible for through our programs" (Michael Williams. "[With Misinformation Swirling in Hurricane Helene's Wake, Officials Urge Residents to 'Stop This Conspiracy Theory Junk'.](#)" CNN. October 4, 2024).

AI-generated images of children fleeing Hurricane Helene, as well as old clips of different storms and computer-generated videos were spreading false information and fueling conspiracy theories about the hurricane and about the government's ability to geo-engineer the weather. Much of the false information about Hurricane Helene and Hurricane Milton was spread on X, which has rules against faked AI content and "Community Notes" for added context but had removed a

feature that allowed users to report misleading information. According to the Institute of Strategic Dialogue, fewer than three dozen false or abusive posts were viewed 160 million times on X, illustrating that misinformation was going viral and being spread to more people than during previous storms (Marianna Spring. " [How Hurricane Conspiracy Theories Took over Social Media](#)." *BBC*. October 10, 2024).

Even after AI-generated images were flagged, some people doubled down on the fake information. For example, a Republican National Committee member representing Georgia responded on X, regarding an image of a young survivor of Hurricane Helene, that she didn't know where the photo came from but that it didn't matter. Other fake images spread online included AI-generated images of animals on rooftops after Hurricanes Helene and Milton. Some experts pointed out that it was often harder for people to fact-check images than text, and that hyper-realistic, uncanny AI-generated images could impact viewers even when they knew, on closer inspection, that the images were false (Huo Jingnan. " [AI-Generated Images Have Become a New Form of Propaganda this Election Season](#)." *NPR*. October 18, 2024).

Tech Industry Response to Misinformation and Disinformation

In the past several years, many companies in the tech industry have taken steps to combat the issue of false news dissemination. Alphabet, for example, is employing third-party fact-checking services to display information about claims regarding controversial issues. In the run-up to the 2016 U.S. presidential election, the firm employed a similar system in its Google News service; after it was deemed to be successful, Alphabet expanded the service in a number of other countries, before finally offering the service on a site-wide basis (Nick Summers. " [Google Will Flag Fake News Stories in Search Results](#)." *Engadget*. April 7, 2017). Further, since 2016, Alphabet has funded almost 20 European projects aimed at fact-checking potentially false reports (Mark Scott. " [In Europe's Election Season, Tech Vies to Fight Fake News](#)." *New York Times*. May 1, 2017). In April 2017, the firm [responded](#) to the phenomenon by enhancing its search ranking system and adding direct feedback tools to its search function so that misleading information can be flagged by users. In March 2018, Alphabet [announced](#) the launch of its Google News Initiative, a program in which it committed to investing \$300 million over the next three years to improve the accuracy and quality of news appearing on its platforms. [Alphabet](#), [Twitter](#), and Meta also committed in November 2017 to using "trust indicators" to assist users in determining the reliability of news articles (Seth Fiegerman. " [Facebook, Google, Twitter to Fight Fake News with 'Trust Indicators'](#)." *CNN*. November 16, 2017). Additionally, both Alphabet and Meta were supporters of CrossCheck, a collaborative journalism verification project focused primarily on validating online material related to the 2017 French presidential election (" [Fake News: Facebook and Google Team Up with French Media](#)." *BBC News*. February 6, 2017).

Meta also pursued initiatives to address this issue following the 2016 presidential election. For example, it published a blog [post](#) outlining some of the technical changes it made to address the propagation of fake news. These include ways for users to report content as 'fake news,' which could then lead to a post being flagged upon appearing in a user's News Feed. Additionally, Meta stated that it had started a program to work with a third-party fact-checking organization, that it would review articles that people were not sharing after reading in an attempt to identify potentially misleading stories, and that it would work to disrupt the financial incentives for spammers (Gina Hall. " [Facebook Promotes Former New York Times Staffer in Effort to Thwart Fake News](#)." *Silicon Valley Business Journal*. May 1, 2017).

In early March 2024, Alphabet [updated](#) its algorithm to reduce low-quality spam content in Search results, specifically to identify and penalize low-quality AI-generated content falling outside of the firm's quality standards. Alphabet could be applying machine learning to detect poor-quality AI content, such as indicated by irregularities in syntax, lack of nuance, unnatural language, or illogical or abrupt shifts in content flow. A recent study of Google Search results found that the firm's algorithm does favor content created by humans over AI-generated content, with human-created content taking 83% of the top five search results (Samantha Dunn. " [Google Algorithm Update Focuses on Quality, Targets AI-Generated Spam Content](#)." *CCN*. May 2, 2024).

OpenAI announced plans to combat election misinformation in early 2024. The firm released a blog post reviewing its existing policies to fight misuse of generative AI tools, as well as discussing new initiatives. Among the steps it laid out, OpenAI stated that it would ban users from creating chatbots that impersonate real candidates or governments in order to discourage people from voting or misrepresent how voting works. Further, it explained that it won't allow users to build applications for political campaigning and lobbying until more research can be done on its technology. OpenAI also began permanently marking AI images created using its DALL-E image generator with information about their origin. Additionally, OpenAI stated that it was partnering with the National Association of Secretaries of State to direct ChatGPT users asking about voting logistics to accurate sources of information through the nonpartisan website CanIVote.org (Ali Swenson. " [Here's How ChatGPT Maker OpenAI Plans to Deter Election Misinformation in 2024](#)." *AP News*. January 16, 2024).

Ahead of the 2024 EU elections, Alphabet [announced](#) that Google's Jigsaw unit would run a series of animated ads across platforms in Belgium, France, Germany, Italy, and Poland featuring "prebunking" techniques to help viewers identify manipulative content before watching it. The ads were developed in partnership with researchers at the Universities of Cambridge and Bristol, and will be translated into all 24 official EU languages. The head of research at

Jigsaw stated that it was a strategy that worked equally effectively across the political spectrum. Results, including survey responses and the number of people reached, are expected to be published in the summer of 2024 (Martin Coulter. ["Exclusive: Google to Launch Anti-Misinformation Campaign Ahead of EU Elections."](#) Reuters. February 16, 2024).

In April 2024, Alphabet's Google Ventures invested \$20 million in the AI-powered social media monitoring platform Alethea Group. The startup planned to use the funding to build its product to detect online foreign disinformation campaigns. Alethea's machine-learning platform gathers news stories, social media posts, and other content to track adverse user behaviors, the proliferation of new accounts, and other indicators of disinformation. A recent report from Alethea suggested that Russian operatives were creating fake versions of real news websites to link to in their social media campaigns, a significantly more sophisticated approach than used in 2016 (Sam Sabin. ["Exclusive: AI Disinfo Detection Startup Raises \\$20M."](#) Axios. April 9, 2024).

OpenAI shared a blog post in May 2024 explaining how its chatbots were used to spread political spam on the messaging app Telegram channels and Russian propaganda on X, in addition to generating entire articles posted online. It then reported on how the firm was working to eliminate such misuse. OpenAI also banned a group of accounts linked to Spamouflage, which the Company and government agencies had linked to the Chinese government. In addition to its own monitoring, OpenAI affirmed that it was sharing information about these kinds of misinformation campaigns with other AI companies and researchers (Chase DiFelicianantonio. ["OpenAI Says ChatGPT Is Already Being Used to Spread Misinformation. Here's How."](#) San Francisco Chronicle. May 30, 2024).

Yet another example of OpenAI's efforts to confront AI misinformation came in August 2024, when the firm stated that it identified and banned several accounts connected to an Iranian influence campaign, Storm-2035, using generative AI to spread misinformation, including regarding the U.S. presidential election. At the time, OpenAI asserted that the Iranian effort did not seem to reach a sizable audience (Cade Metz. ["OpenAI Says It Disrupted an Iranian Misinformation Campaign."](#) The New York Times. August 23, 2024).

Leaks

In 2023, a researcher who was granted access by Meta to LLaMA (a large language model built by the tech company) leaked it for anyone to download online. This immediately raised concerns that this technology could be used to manufacture of large amounts of fake news, spam, phishing emails, disinformation, and incitement. In response, Meta stated that "[t]here is still more research that needs to be done to address the risks of bias, toxic comments, and hallucinations in large language models" and that "[l]ike other models, LLaMA shares these challenges." Meta also stated that, consistent with previous models, LLaMA was shared for research purposes. Further, Meta stated that while the model is not accessible to all, and some have tried to circumvent the approval process, the firm believes the current release strategy allows it to balance responsibility and openness (Katyanna Quach. ["LLaMA Drama as Meta's Mega Language Model Leaks."](#) The Register. March 8, 2023).

TECH INDUSTRY CONCERNS REGARDING USE OF AI

Alongside the concerns raised by legislators and consumers, tech industry leaders have also voiced concerns about the direction of travel of AI technology. In March 2023, many of these leaders signed onto a letter calling for a six-month pause to work on systems "more powerful" than the GPT-4 framework. Apple's co-founder was among the letter's 1,800 signatories, which also included Elon Musk, cognitive scientist Gary Marcus, and engineers from the Company, Amazon, DeepMind, Meta, and Alphabet. The release of the letter also sparked backlash from some researchers who claimed that their signatures were added without permission, as well as from others whose work was cited in the letter. They criticized the letter's focus on imagined apocalyptic scenarios over the more immediate risks of the programming of racist or sexist bias into machines (Kari Paul. ["Letter Signed By Elon Musk Demanding AI Research Pause Sparks Controversy."](#) The Guardian. April 1, 2023).

Apple's co-founder further expressed concerns regarding AI. In May 2023, he explained that AI could be harnessed by "bad actors" and could make it harder for people to spot scams and misinformation. Further, he stated that responsibility for AI-generated content should rest with the people who publish that content and that the sector needed regulation (Philippa Wain. ["Apple Co-Founder Says AI May Make Scams Harder to Spot."](#) BBC News. May 8, 2023).

REGULATIONS GOVERNING AI

United States

U.S. lawmakers and regulators have expressed increasing concern over the potential risks of generative AI as such models become more advanced and widespread. For instance, in a May 2023 meeting at the White House, President Joe Biden and Vice President Kamala Harris pushed AI developers, including Alphabet, OpenAI, and the Company, to seriously consider concerns over the use of AI. They also pushed for developers to be more open about their products, the need for AI systems to be subjected to outside scrutiny, and the importance that those products be kept away from bad actors. The White House further pledged to release draft guidelines for government agencies' use of AI safeguards

(David McCabe. " [White House Pushes Tech C.E.O.s to Limit Risks of A.I.](#)" *The New York Times*. May 4, 2023). Following the meeting, the Company, OpenAI, and five other companies engaged in AI development agreed to voluntary safeguards. The safeguards include security testing, in part by independent experts; research on bias and privacy concerns; information sharing about risks with governments and other organizations; development of tools to fight societal challenges like climate change; and transparency measures to identify AI-generated material (Michael D. Shear, Cecilia Kang, David E. Sanger. "[Pressured by Biden, A.I. Companies Agree to Guardrails on New Tools](#)." *The New York Times*. July 21, 2023).

Despite these voluntary measures, Biden signed an [executive order](#) in October 2023 requiring AI developers to share their safety test results and other information with the government. The order also directed government agencies to create new standards to ensure AI tools are safe and secure before public release, required new guidance to label and watermark AI-generated content to help differentiate them from authentic interactions, and asked federal agencies to review the use of AI in the criminal justice system (Josh Boak, Matt O'Brien. "[Biden Wants to Move Fast on AI Safeguards and Signs an Executive Order to Address His Concerns](#)." *Associated Press*. October 30, 2023).

A year after Biden's executive order on AI, the White House [published](#) a press release providing updates on the order's initiatives. It stated that Federal agencies had completed on schedule each action that the order tasked for the past year, more than one hundred in total. For instance, agencies used the Defense Production Act to require developers of the most powerful AI systems to report vital information, including on safety and security testing, to the U.S. government. Further, the U.S. AI Safety Institute ("U.S. AISI") at the Department of Commerce began pre-deployment testing of major new AI models through recently signed agreements with two leading AI developers. Additionally, the U.S. AISI and the National Institute of Standards and Technology at the Department of Commerce published frameworks for managing risks related to generative AI and dual-use foundation models, and the AISI released a request for information on responsible development and use of AI models for chemical and biological sciences. Beyond other measures to manage risks to safety and security, the administration also reported on initiatives taken to: (i) stand up for workers, consumers, privacy, and civil rights, including releasing the Department of Labor's AI Principles and Best Practices; (ii) harness AI for good, including establishing two new National AI Research Institutes for building AI tools to advance progress across economic sectors, science, and engineering; (iii) bring AI and AI talent into government; and (iv) advance U.S. leadership abroad.

AI developers have also attracted attention from the legislative and judicial branches. In September 2023, U.S. Senators Richard Blumenthal and Josh Hawley announced their plans to introduce a framework to regulate AI technology. The framework included requirements for the licensing and auditing of AI, the creation of an independent federal office to oversee the technology, liability for companies for privacy and civil rights violations, and requirements for data transparency and safety standards. The Senate also met with industry leaders, including from the Company and OpenAI, in a separate meeting on possible AI-related regulations (Cecilia Kang. " [2 Senators Propose Bipartisan Framework for A.I. Laws](#)." *The New York Times*. September 7, 2023). Meanwhile, in a September 2023 letter, the attorney generals of all 50 U.S. states urged Congress to study how AI could be used in child exploitation. The letter further called on Congress to expand existing restrictions on child sexual abuse materials specifically to cover AI-generated images (Meg Kinnard. "[Prosecutors in All 50 States Urge Congress to Strengthen Tools to Fight AI Child Sexual Abuse Images](#)." *Associated Press*. September 5, 2023).

Around the same time, U.S. Senators Ron Wyden and Cory Booker, along with Representative Yvette Clarke, [introduced](#) the [Algorithmic Accountability Act of 2023](#), to create new protections for people affected by AI systems that are impacting decisions affecting credit, housing, education and other high-impact uses. The bill applies to new generative AI systems used for critical decisions, as well as other AI and automated systems, and would obligate the FTC to require companies to perform impact assessments of their AI systems. It would also create a public repository at the FTC of these systems.

Also in September 2023, Senator Amy Klobuchar introduced S.2770, [Protect Elections from Deceptive AI Act](#). This bill would prohibit individuals, political committees, and other entities from knowingly distributing materially deceptive AI-generated audio or visual media of a federal candidate, or in carrying out a federal election activity, with the intent of influencing an election or soliciting funds. However, this prohibition would not apply to radio or television broadcasting stations that broadcast such media with disclosures as part of a bona fide newscast. In addition, the bill would permit a federal candidate whose voice or likeness appears in materially deceptive AI-generated audio or visual media to bring civil action for injunctive relief or damages.

There have also been a number of state-level regulations introduced governing AI. For example, as of June 2023, nine states had enacted regulations relating to deepfake content, most often in the context of pornography and election influence, and four other states were pursuing similar bills. The first states to pass deep fake legislation include California, Texas, and Virginia, who passed bills in 2019, while Minnesota passed its deepfake law in May 2023 and Illinois is waiting for its governor to sign its new deepfake legislation. Many states are also amending their election codes to ban deepfake campaign ads within a specific time frame before an election. Yet another complication to creating legislation to address harmful deepfakes is that some experts have expressed concern that well-meaning legislation, if not carefully crafted, could have a detrimental effect on people's First Amendment rights. For example, the ACLU of Illinois at first

opposed the pornographic deepfake bill developed in Illinois because the bill's sweeping provisions and immediate takedown clause could "chill or silence vast amounts of protected speech." In response, lawmakers have employed various amendments to change the bill to include deepfakes into the state's existing revenge porn statute, which the ACLU said was an improvement, though it still maintained some concern. California, on the other hand, included specific references to First Amendment protections in its bill. Any federal regulations on AI-generated deepfakes will likely face the same concerns regarding free speech, especially if they include broad language such as limiting exceptions to "legitimate public concern" (Isaiah Poritz. "[Deepfake Porn, Political Ads Push States to Curb Rampant AI Use](#)." *Bloomberg Law*. June 20, 2023). In the 2024 legislative session, at least 45 states, Puerto Rico, the Virgin Islands, and Washington D.C., [introduced](#) AI-related bills and 31 states, Puerto Rico, and the Virgin Islands adopted resolutions or enacted legislation on this topic.

The European Artificial Intelligence Act

After months of negotiations between different political groups, European lawmakers agreed on the EU AI Act, the world's first comprehensive set of rules governing AI technology, in May 2024 (Martin Coulter. "[Tech Giants Push to Dilute Europe's AI Act](#)." *Reuters*. September 20, 2024). The Act came into [effect](#) on August 1, 2024, introducing a framework across all EU countries, based on a forward-looking definition of AI and a risk-based approach:

- Minimal risk: most AI systems such as spam filters and AI-enabled video games face no obligation under the Act, but companies can voluntarily adopt additional codes of conduct;
- Specific transparency risk: systems like chatbots must clearly inform users that they are interacting with a machine, while certain AI-generated content must be labeled as such;
- High risk: high-risk AI systems such as AI-based medical software or AI systems used for recruitment must comply with strict requirements, including risk-mitigation systems, high-quality of data sets, clear user information, human oversight, etc.; and
- Unacceptable risk: for example, AI systems that allow "social scoring" by governments or companies are considered a clear threat to people's fundamental rights and are therefore banned.

Nevertheless, lawmakers still had to determine the accompanying codes of practice and the EU invited companies, researchers, and others to help in the drafting. The EU received almost 1,000 applications, with the largest tech companies encouraging a light-touch approach to the law's implementation. The code of practice would come into effect in late 2025 but would not be legally binding; rather, it would provide a compliance checklist to companies. The EU AI Act addresses a variety of AI issues, including data scraping practices that could potentially breach copyright. Companies would have to disclose "detailed summaries" of the data used to train models, potentially providing an avenue for legal recourse to creators whose work was used to train AI models. Some business leaders worried that the Act's requirements could render companies' trade secrets vulnerable and criticized the EU for prioritizing regulation over innovation (Martin Coulter. "[Tech Giants Push to Dilute Europe's AI Act](#)." *Reuters*. September 20, 2024).

A new tool from Swiss startup LatticeFlow AI and partners ETH Zurich and Bulgaria's INSAIT tested generative-AI models like those from OpenAI and Meta, among others, across dozens of categories and in line with the new EU AI Act, which comes into effect in stages over the next two years. While the results showed that models developed by OpenAI, Meta, Mistral, Alibaba, and Anthropic all received average scores of 0.75 or above (out of scores between 0 and 1), some models demonstrated shortcomings in key areas. Regarding discriminatory output, OpenAI's GPT-3.5 Turbo scored a relatively low 0.46, and Alibaba Cloud's Qwen1.5 72B Chat scored only 0.37. When tested for prompt hijacking, a type of cyberattack using malicious prompts, Meta's Llama 2 13B Chat model scored 0.42 and Mistral's 8x7B Instruct scored 0.38. The LatticeFlow LLM checker would be extended to include further measures as they were introduced, and LatticeFlow stated that it would be freely available for developers to test their models' compliance online. The European Commission said it welcomed the study and evaluation platform as a first step in translating the EU AI Act into technical requirements. Companies failing to comply with the EU AI Act could face fines of 35 million euros or \$38 million (Martin Coulter. "[Exclusive: EU AI Act Checker Reveals Big Tech's Compliance Pitfalls](#)." *Reuters*. October 16, 2024).

European Code of Practice on Disinformation

In response to this proposal, the Company states that it publishes a broad set of reports on its efforts to address misinformation and disinformation, and that it has taken steps to meet its obligations as a signatory to the European Code of Practice on Disinformation. The European Commission created the EU Code of Practice on Disinformation ("Code") in 2018, which at the time was the first self-regulatory piece of legislation that intended to prompt companies to collaborate on solutions to the problem of disinformation. After an initial amendment process, the EU invited companies to sign the new strengthened code in 2022, including large tech companies such as the Company, Meta, and Google. Notably absent from the list of [signatories](#) are Amazon, Apple, and Telegram. Additionally, some companies decided to select which measures they would comply with. For example, Meta, Google, TikTok, and Twitter did not agree to implement "trustworthiness indicators" to inform users about possible disinformation on their sites (Brooke Tanner. "[EU Code of Practice on Disinformation](#)." Brookings. August 5, 2022).

The [2018 Code of Practice on Disinformation](#) made a commitment to "deploy policies and processes to disrupt advertising and monetization incentives for relevant behaviors, such as misrepresenting material information about oneself or the purpose of one's properties." However, because this was a broad definition of disinformation and did not require signatories to commit to any specific measures, the code was revised in 2022. The [2022 Strengthened Code of Practice on Disinformation](#) empowered platforms and industry to adhere to self-regulatory standards to combat disinformation. For example, the new code [empowered](#) companies to contribute to wide-ranging improvements by signing up to precise commitments relevant to their field, such as demonetising the dissemination of disinformation; guaranteeing transparency of political advertising; enhancing cooperation with fact-checkers; and facilitating researchers access to data.

Australian Code of Practice for Disinformation and Misinformation

The Company affirms in its response to this proposal that it files reports under the [Australian Code of Practice for Disinformation and Misinformation](#). Launched in February 2021 by the Digital Industry Group, a non-profit association advocating for the digital industry in Australia, the Australian Code of Practice on Disinformation and Misinformation commits technology companies to reducing the risk of online misinformation causing harm to Australians. The code has been adopted by the Company, Adobe, Apple, Facebook, Google, TikTok, and others. All signatories commit to safeguards to protect Australians against harm from online disinformation and misinformation, and to adopting measures that reduce its spread and visibility. Participating companies also commit to releasing an annual transparency report about their efforts under the code, which will help improve understanding of online misinformation and disinformation in Australia.

The Australian Communications and Media Authority ("ACMA") oversees the code of practice. In addition to submitting annual progress reports to the ACMA, signatories [commit](#) to:

- Enacting safeguards against harm arising from the spread of misinformation and disinformation on their platforms;
- Disrupting advertising and monetisation incentives for the spread of misinformation and disinformation;
- Working to ensure the security and integrity of the platform's services and products;
- Empowering users to make better-informed choices of digital content and helping them identify false and misleading content;
- Increasing transparency around political advertising;
- Supporting research that improves public understanding of misinformation and disinformation; and
- Publicising the measures platforms are taking to combat misinformation and disinformation.

Other International Regulations

Meanwhile, both Israel and Japan quickly clarified their existing regulations pertaining to data, privacy, and copyright protections, both with language that enabled AI to train with copyrighted content. Responding more broadly, the United Arab Emirates has developed draft legislation, working groups on AI best practices, and sweeping proclamations regarding AI strategy. Still many other countries are opting for a wait-and-see approach, despite the array of warnings regarding the need for international cooperation on AI regulation and inspection, including a statement from OpenAI's CEO in May 2023 emphasizing the "existential risk" of AI technology (Mikhail Klimentov. " [From China to Brazil, Here's How AI Is Regulated Around the World](#)." *The Washington Post*. September 3, 2023).

For example, in Brazil, legislators have developed a 900-page draft of a bill that would require companies to conduct risk assessments before bringing a product to market, and it explicitly bans AI systems that deploy "subliminal" techniques or exploit users in ways that are harmful to their health or safety. It would also require a government database to publicize AI products determined to have "high risk" implementations and AI developers would be liable for damage caused by their AI systems. Similar to the EU AI Act, the Brazilian draft legislation categorizes different types of AI based on the risk they pose to society (Mikhail Klimentov. " [From China to Brazil, Here's How AI Is Regulated Around the World](#)." *The Washington Post*. September 3, 2023).

There have also been attempts to regulate AI within the UK. In March 2023, the UK government published an [AI policy paper](#) that outlined proposals for regulating the use of AI within the country, though it stated that it will refrain from regulating the British AI sector "in the short term" (Daria Mosolova. "[UK Will Refrain from Regulating AI 'In the Short Term'](#)." *Financial Times*. November 16, 2023). In the long-term, however, the UK government has signaled its [intent](#) to draft a central cross-economy AI risk register and update its AI regulation roadmap with indications of whether a government unit or independent body would be the most appropriate mechanism to deliver the central functions.

There is also potential regulatory momentum in China. In January 2024, the country issued draft guidelines for standardizing the AI industry, proposing the formation of more than 50 national and industry-wide standards by 2026 and stating its intention to participate in the formation of more than 20 international standards by the same time. China's industry ministry stated that the aim of 60% of the prospective standards should be serving "general key technologies and application development projects," and it has targeted more than 1,000 companies to adopt and advocate for the new standards (Josh Ye. "[China Issues Drafts Guidelines for Standardizing AI Industry](#)." *Reuters*. January 18, 2024). China

has also created a draft bill regarding generative AI, with a translation indicating that developers would "bear responsibility" for outcry created by their AI. The draft bill would also hold developers liable if their use of training data infringed on someone else's intellectual property, and it also states that the design of AI services must lead to the generation of only "true and accurate" content (Mikhail Klimentov. " [From China to Brazil, Here's How AI Is Regulated Around the World](#)." *The Washington Post*. September 3, 2023).

There have also been multilateral attempts at regulating AI technologies. In May 2023, the U.S.-EU Trade and Technology Council [stated](#) in May 2023 their intention to develop a voluntary AI Code of Conduct as concerns grow about the risks AI poses to humanity (" [US, Europe Working on Voluntary AI Code of Conduct As Calls Grow For Regulation](#)." *AP News*. May 31, 2023). In late October 2023, leaders of the Group of Seven ("G7") economies (made up of Canada, France, Germany, Italy, Japan, Britain, the U.S., and the EU), agreed to the voluntary [code of conduct](#). The code includes 11 points and aims "to promote safe, secure, and trustworthy AI worldwide and provides voluntary guidance for actions by organizations developing the most advanced AI systems, including the most advanced foundation models and generative AI systems". The code urges companies to take appropriate measures to identify, evaluate, and mitigate risks throughout the AI life-cycle and to address incidents and patterns of misuse after AI products have been released on the market (Foo Yun Chee. " [Exclusive: G7 to Agree AI Code of Conduct For Companies](#)." *Reuters*. October 29, 2023).

COMPANY AND PEER DISCLOSURE

Company Disclosure

The Company acknowledges that it is investing in AI across the entire Company and infusing generative AI capabilities into its consumer and commercial offerings. It notes that it expects AI technology and services to be a highly competitive and rapidly evolving market, as new competitors continue to enter the market. The Company states that it will bear significant development and operational costs to build and support the AI models, services, platforms, and infrastructure necessary to meet the needs of its consumers. Moreover, it explains that to compete effectively, it must be responsive to technological change, new and potential regulatory developments, and public scrutiny (2024 10-K, p.21). The Company notes that potential risks due to issues in the development and use of AI may result in reputational or competitive harm or liability. It explains that it is building AI into many of its offerings and is also making it available for its customers to use in solutions that they build, and this AI may be developed by the Company or others, including its strategic partner, OpenAI. It asserts that, as with many innovations, AI presents risk and challenges that could affect its adoption, and therefore its business. It states that AI algorithms or training methodologies may be flawed; datasets may be overboard, insufficient, or contain biased information; and content generated by AI systems may be offensive, illegal, inaccurate, or otherwise harmful. Moreover, it adds that ineffective or inadequate AI development or deployment practices by the Company or others could result in incidents that impair the acceptance of AI solutions; cause harm to individuals, customers, or society; or result in the Company's products and services not working as intended. It also states that human review of certain outputs may be required. The Company states that its implementation of AI systems could result in legal liability; regulatory action; brand, reputational, or competitive harm; or adverse impacts. It explains that these risks may arise from current copyright infringement and other claims related to AI training and output, new and proposed legislation and regulations, such as the EU AI Act and the U.S. AI Executive Order, and new applications of data protection, privacy, consumer protection, intellectual property, and other laws. Further, it explains that some AI scenarios present ethical issues or may have broad impacts on society, and if it enables or offers AI solutions that have unintended consequences; unintended usage or customization by its customers and partners; are contrary to the Company's responsible AI policies and practices; or are controversial because of their impact on human rights, privacy, employment, or other social, economic, or political issues, then the Company's reputation, competitive position, business, financial condition, and results of operations may be adversely affected (2024 10-K, p.26).

The Company also discusses cyber influence operations and other cybersecurity and cybercrime issues in its [Digital Defense Report](#). Specifically, the Company details the risk of AI's impact on cybersecurity and specifically the emerging threat landscape, including:

- Impersonation (the use of image, audio, and video deepfakes to impersonate individuals, with specific threats including fraud, defamation, and information warfare);
- Content production (the creation of harmful content for dissemination such as disinformation);
- Nefarious knowledge acquisition;
- Cyber threat amplification (automated generation of malware and attack command and control infrastructure);
- Direct social attacks (malicious activities such as propaganda, phishing, and terrorist recruitment); and
- Indirect social attacks (automated harassment and defamation), among other things.

(pp.87-88)

The Company [states](#) that in the coming year, it anticipates the biggest rises in automated fraud and election interference, CSAM and non-consensual intimate image production, and the use of XPIA and deepfake impersonation as cyberattack and fraud channels (p.87). It also examines sophisticated AI-enabled human targeting, as well as emerging techniques in

AI-enabled attacks, including deepfakes and other variations on social engineering. Despite the increasing risks, the Company states that by incorporating AI into risk mitigation activities, defenders can evolve at the same or a greater rate as threat actors (pp.89-90). The Company also discusses nation-state threat actors using AI for influence operations, and it charts adversarial use of AI in influence operations from China, Russia, and Iran and proxies (pp.91-92). Regarding limiting foreign influence operations in the modern era, the Company emphasizes the need for international norms to regulate foreign influence operations in the online space. It recommends that states embrace limitations on foreign influence operations, including limits on targets regarding crisis/emergency scenarios, emergency/humanitarian response organizations, elections (stating that covert interference in elections via foreign influence operations must be prohibited), and vulnerable/marginalized communities. Further, it recommends limits on tools and techniques regarding the covert use of AI and theft/abuse of social media data (p.93).

The Company [addresses](#) how it helps support democratic elections, explaining on page 79 how its initiatives stem from four key principles:

1. Voters have a right to transparent and authoritative information regarding elections;
2. Candidates should have the ability to verify the authenticity of content originating from their campaigns and have access to procedural or legal mechanisms to address instances where their likeness or content is manipulated by AI to mislead the public during elections;
3. Political campaigns should have the resources to safeguard against cyber threats and effectively utilize AI, with access to affordable, easily deplorable tools, training, and support; and
4. Election authorities should be able to ensure a secure and resilient election process and have access to tools and services that enable this process.

The Company [reviews](#) its steps for detection, response and mitigation, as well as collaboration. It lays out how it is helping protect the online environment surrounding elections by defending the information environment, protecting data, identifying and responding to threats, and boosting public awareness of AI elections risk (pp.79-80).

The Company also [provides](#) an [update](#) on its AI safety policies that outlines its approach to advancing responsible AI, including by implementing voluntary commitments that it and others made at the White House convening in July 2023, and discusses its collaboration with OpenAI in this area. The update discusses in detail the Company's approach in nine areas of practice and investment:

- Alignment of its efforts with the White House Voluntary AI Commitments;
- Responsible capability scaling;
- Model evaluations and red teaming;
- Model reporting and information sharing;
- Security controls, including securing model weights;
- Reporting structure for vulnerabilities found after model release;
- Identifiers of AI-generated material;
- Prioritizing research on risks posed by AI;
- Preventing and monitoring model misuse; and
- Data input controls and audit.

The Company's [responsible AI](#) webpage features its latest news and announcements on responsible AI policy, research, and engineering, as well as the Company's tools and policy for advancing AI. It also identifies the Company's six principles for AI development and use, including:

- Fairness: AI systems should treat all people fairly;
- Reliability and safety: AI systems should perform reliably and safely;
- Privacy and security: AI systems should be secure and respect privacy;
- Inclusiveness: AI systems should empower everyone and engage people;
- Transparency: AI systems should be understandable; and
- Accountability: People should be accountable for AI systems.

Additionally, the Company [provides](#) details on its Information Integrity Principles, which include:

- Freedom of expression;
- Authoritative content;
- Demonetization; and
- Proactive efforts.

The Company also discloses its first [Responsible AI Transparency Report](#), addressing: (i) how the Company builds generative applications responsibly; (ii) how it makes decisions about releasing generative applications; (iii) how it supports its customers in building responsibly; and (iv) how it learns, evolves, and grows, which includes a discussion of how industry and others stand to significantly benefit from the key role that governments can play when using

consensus-based safety frameworks (pp.4-32).

Regarding disinformation, the Company [states](#) that amid growing concern that AI can make it easier to create and share disinformation, the Company recognizes that it is imperative to give users a trusted experience, but adding that as generative AI technologies become more advanced and prevalent, it is increasingly difficult to identify AI-generated content. The Company explains that labeling AI-generated content and disclosing when and how it was made (otherwise known as provenance) is one way to address this issue, and in May 2023, the Company announced its intent to build new media provenance capabilities that use cryptographic methods to mark and sign AI-generated content with metadata about its source and history. It further states that the Company has made significant progress on its commitment to deploy new state-of-the-art tools to help the public identify AI-generated audio and visual content, and by the end of 2023, it was automatically attaching provenance metadata to images generated with OpenAI's DALL-E 3 model in its Azure OpenAI Service, Microsoft Designer, and Microsoft Paint (p.13). The Company explains that its provenance metadata, referred to as Content Credentials, is applied by using an open technical standard developed by the [Coalition for Content Provenance and Authenticity](#) ("C2PA"), which the Company co-founded in 2021. It [notes](#) that there are now more than 100 industry members of C2PA, and in February 2024, OpenAI announced that it would implement the C2PA standard for images generated by its DALL-E 3 image model (p.13).

The Company also [explains](#) that to address the risk posed by bad actors' use of deceptive AI content to influence elections in 2024, it worked with 19 other companies, including OpenAI, to announce the new Tech Accord to Combat Deceptive Use of AI in 2024 Elections at the Munich Security Conference in February. It notes that the commitments of the Accord include advancing provenance technologies, innovating robust disclosure solutions, detecting and responding to deepfakes in elections, and fostering public awareness and resilience. It adds that since signing the Tech Accord, the Company has launched a portal for candidates to report deepfakes on its services, and in March, it launched Microsoft Content Integrity tools in private preview, to help political candidates, campaigns, and elections organizations maintain greater control over their content and likeness. It states that the Content Integrity tools include two components: one to certify digital content by adding Content Credentials, and another to allow the public to check if a piece of digital content has Content Credentials (p.14). Moreover, the Company discusses its [Democracy Forward](#) initiative, which [provides](#) journalists and newsrooms with tools and technology to help build capacity, expand their reach and efficiency, distribute trustworthy content, and ultimately provide the information needed to sustain healthy democracies. Further, the Company discusses its partnerships with the National Association of State Election Directors, as well as its endorsement of the Protect Elections from Deceptive AI Act in the U.S., and its Elections Communications Hub (p.14).

In its response to this proposal, the Company notes that every six months, it publicly reports on the steps it has taken to meet its obligations as a signatory to the European Code of Practice on Disinformation, and it emphasizes that these [reports](#) include detailed quantitative metrics on a range of measures related to disinformation. It states that it also files similar reports under the Australian Code of Practice for Disinformation and Misinformation, as well. The Company explains that it established the Microsoft Threat Analysis Center ("MTAC"), bringing together analysts focused on nation-state driven cyber and influence actors to identify, analyze, and expose cyber-enabled influence operations targeting democracies around the world. It states that MTAC has [issued](#) regular reporting on threats to elections around the world (2024 DEF 14A, pp.85-86). In November 2023, MTAC [released](#) a report, "[Protecting Election 2024 from Foreign Malign Influence](#)," and it [discussed](#) grounding the Company's Election Protection Commitments in a set of principles to help safeguard voters, candidates, campaigns, and elections authorities worldwide.

In May 2023, the Company published [Governing AI: A Blueprint for the Future](#), which examines the Company's legal and regulatory blueprint for the future and discusses its approach to building AI systems that benefit society. The first portion of the document considers a legal and regulatory blueprint for the future, including:

- Implementing and building upon new government-led AI safety frameworks;
- Requiring effective safety brakes for AI systems that control critical infrastructure;
- Developing a broad legal and regulatory framework based on the technology architecture for AI;
- Promoting transparency and ensuring academic and nonprofit access to AI; and
- Pursuing new public-private partnerships to use AI as an effective tool to address the inevitable societal challenges that come with new technology.

The second portion of the Blueprint [examines](#) the Company's approach to responsible design and to building AI systems that benefit society, including:

- Committing to developing AI responsibly;
- Operationalizing responsible AI at the Company;
- Applying the Company's Responsible AI approach to the new Bing;
- Advancing of Responsible AI through Company culture; and
- Empowering customers on their Responsible AI journey.

Additionally, the Company provides its most recent [Responsible AI Standard](#), which defines the Company's product

development requirements and goals, aligned with its responsible AI principles. As one of its transparency goals, the Company states that its AI systems are designed to inform people that they are interacting with an AI system or are using a system that generates or manipulates image, audio, or video content that could falsely appear to be authentic (p.12). Additionally, regarding reliability and safety, the Company states that it designs its AI systems to minimize the time to remediation of predictable or known failures, as well as that AI systems are subject to ongoing monitoring, feedback, and evaluation to identify and review new uses, identify and troubleshoot issues, manage and maintain the systems, and improve them over time (pp.23-25). The Company also provides a [Responsible AI Impact Assessment Template](#) and its [Responsible AI Standard Reference Guide](#). Its website on responsible AI [principles and approach](#) also examines responsible AI in action as well as how the Company implements responsible AI and other transparency notes.

The Company also maintains a set of [AI Customer Commitments](#) and an [AI learning hub](#). Further, it [provides](#) a learning module on responsible and trusted AI principles, along with an [introduction](#) to its guidelines for human-AI interaction.

Further speaking to the risks of its AI systems, OpenAI provides a May 2023 statement presenting its specific concerns regarding the [governance of superintelligence](#), as well as discussions regarding [developing safe and responsible AI](#), its [approach to alignment research](#) on AI, its [approach to AI safety](#), questions of [how AI should behave and who should decide](#) and also of [democratic inputs to AI](#), specifically looking at funding experiments to set up a democratic process for deciding what rules AI systems should follow, within the bounds defined by law, among [other topics](#). In reference to governance, OpenAI [states](#) that it will likely be necessary to have something comparable to the International Atomic Energy Agency but for superintelligence efforts and that any effort above a certain capability (or resources like compute) threshold will need to be subject to an international authority that can inspect systems, require audits, test for compliance with safety standards, place restrictions on degrees of deployment and levels of security, etc. Continuing, OpenAI explains that tracking compute and energy usage could go a long way and give some hope this idea could actually be implementable. It then adds that as a first step, companies could voluntarily agree to begin implementing elements of what such an agency might one day require, and as a second, individual countries could implement it. The statement stresses that it would be important that such an agency focus on reducing existential risk and not issues that should be left to individual countries, such as defining what an AI should be allowed to say. It also broaches the topic of needing the technical capability to make a superintelligence safe.

In early 2024, OpenAI [published](#) information regarding how it would be approaching the 2024 worldwide elections, with details on preventing abuse, transparency around AI-generated content, and improving access to authoritative voting information. It stated that it regularly refines its [usage policies](#) for ChatGPT and the API as it [learns](#) more about how people use or attempt to abuse its technology. Further, OpenAI provides updates to its election initiatives for May and October 2024. Specifically regarding election results starting on November 5, OpenAI explains that people who ask ChatGPT for elections results will see a message encouraging them to check news sources like the *Associated Press* and *Reuters*, or their state or local election board for the most complete and up-to-date information.

Regarding oversight of this issue, the [environmental, social, and public policy committee](#) assists the board in overseeing the key non-financial regulatory risks that may have a material impact on the Company and especially its ability to sustain trust with customers, employees, and the public, which includes policies and programs and related risks that concern environmental sustainability, the social and public policy impacts of technology including privacy, digital safety, and responsible artificial intelligence, and legal, regulatory, and compliance matters relating to competition / antitrust, trade, and national security. The committee also reviews and provides guidance to the board and management about key environmental and social matters such as responsible artificial intelligence and human rights, among others.

Peer Analysis

To compare, **Alphabet Inc.** (NASDAQ: GOOGL) discloses its [AI principles](#) along with information on implementing its AI Principles. Additionally, the firm provides several updates it has made since 2019 regarding its AI Principles. It also provides a list of AI applications that it will not pursue. Further, Alphabet [provides](#) information regarding Google Gemini, bringing AI to its platforms, helping developers solve complex challenges, and weaving responsible AI into its work. Moreover, the firm [discusses](#) its responsible approach to building guardrails for generative AI, including its [generative AI prohibited use policy](#).

It also details its [approach](#) to the 2024 elections in the U.S., including safeguarding its platforms from abuse, helping identify AI-generated content, surfacing high-quality information, and partnering to equip campaigns with best-in-class security. The firm requires election advertisers to prominently [disclose](#) when their ads include realistic synthetic content that has been digitally altered or generated, including by AI tools. It also [requires](#) the creators to disclose when they've created altered or synthetic content that is realistic and to display a label indicating such. Alphabet also provides a [Political Advertising Transparency Report](#), as part of its efforts to ensure users that Alphabet's advertisers have passed an [identity verification](#) process.

Alphabet [discusses](#) partnerships with Democracy Works, its [Advanced Protection Program](#), Defending Digital Campaigns, and its Campaign Security Project, among others. It also addresses fighting misinformation online regarding the 2024

elections in [Europe](#), including debunking and prebunking and launching the coalition [Elections24Check](#). In addition, the firm [reviews](#) information on new journalist resources to raise awareness of disinformation, partnering to empower young voters, and providing topical insights on Search trends.

Further, Alphabet provides its YouTube [Community Guidelines](#), [political content](#) policies for advertisers, and also policies on [manipulated content](#), incitement to [violence](#), [hate speech](#), and [harassment](#). Further, it addresses ongoing transparency on [election ads](#) and countering misinformation with its [Google News Initiative Training Network](#), [Fact Check Explorer](#), and Google [Fact Check Tools](#). It also discusses helping people navigate AI-generated content, including [ad disclosures](#), content [labels](#) on YouTube, a responsible approach to generative AI products, providing users with additional [context](#), and digital [watermarking](#).

Additionally, the firm discusses [ongoing work](#) to fight misinformation online and [partnering](#) with news publishers to drive innovation and fight misinformation. Further, it also [presents](#) information on how YouTube addresses misinformation using a [combination](#) of content reviewers and machine learning to identify content that violates its policies. Moreover, Alphabet's quarterly [Transparency Report](#) on Youtube's community guidelines discloses data on policy enforcement and removed channels, discussing how automated flagging systems, along with human review, aid in detecting content that violates company policies.

The firm also provides a third-party [civil rights audit](#), conducted by Wilmer Cutler Pickering Hale and Dorr LLP in 2023, which benchmarked several topics related to AI, including misinformation.

Regarding board oversight, the [audit and compliance committee](#) has responsibility for oversight of risks and exposures associated with data privacy and security, competition, legal, regulatory, compliance, civil and human rights, sustainability, and reputational risks. Alphabet also states in its latest proxy statement that the board's oversight function of major risks and risk exposures, including those relating to or resulting from the firm's development and implementation of AI in its products and services, sits at the top of its risk management framework, and as set forth in Alphabet's Corporate Governance Guidelines, the board is ultimately responsible for covering strategic, financial, and execution risks and exposures associated with the firm's business strategy, production innovation, and policy and significant regulatory matters that may present material risk to its financial performance, operations, plans, prospects, or reputation. Further, it adds that the board's skills and expertise, including deep technical expertise in computer science, facilitates oversight of a highly complex global business, and the full board meetings have regularly and extensively covered AI issues. The audit committee and senior management provide the board with reports and updates regarding issues and risk exposures regarding AI development, and these discussions ensure that the board is fully involved in the oversight of Alphabet's business strategies and plans as they relate to AI (2024 DEF 14A, p.94).

To further compare, **Meta Platforms, Inc.** (NASDAQ: META) states that it has developed tools to help people understand when content has been created using its Meta AI feature, and the firm has ongoing efforts to help identify other AI-generated content on its apps. For example, when photorealistic images are created using its Meta AI feature, Meta does several things to make sure people know AI is involved, including putting visible markers that users can see on the images, and both invisible watermarks and metadata embedded within image files. It adds that using both invisible watermarking and metadata in this way improves both the robustness of these invisible markers and helps other platforms identify them. The firm also states that its goal is to help identify AI-generated content created with other companies' tools as well (2024 DEF 14A, p.84).

Meta [provides](#) its commitment to responsible AI, with details on its five pillars of: (i) privacy and security; (ii) fairness and inclusion; (iii) robustness and safety; (iv) transparency and control; and (v) accountability and governance. It discusses how Meta is actioning responsible AI through datasets, privacy, ad delivery, associations, ad-driven feeds, system cards, and policy. More specifically, the webpage features information on topics such as building diverse datasets and tools for more inclusive AI products, protecting privacy while addressing fairness concerns, innovating to improve fairness in ad delivery, generating responsible AI associations, giving more control over AI-driven feeds and recommendations, developing new methods for explaining the firm's AI systems, and testing new policy approaches to AI transparency, explainability, and governance.

Additionally, the firm provides a [Responsible Use Guide](#) for Meta LLama, and its [Meta Llama Guard](#) approach to trust and safety. The firm further discusses its [approach to AI](#), its Fundamental AI Research [team](#), as well as responsible AI and [social good projects](#). Additionally, it [discusses](#) building generative AI features responsibly.

Meta also provides its [misinformation](#) policy rationale, with several updates between May 2024 and July 2024, stating that Meta removes misinformation where it is likely to directly contribute to the risk of imminent physical harm, and it also removes content that is likely to directly contribute to interference with the functioning of political processes. For misinformation in these categories, the firm partners with independent experts who possess knowledge and expertise to assess the truth of the content and whether it is likely to directly contribute to the risk of imminent harm. It then notes that for all other misinformation, Meta focuses on reducing its prevalence or creating an environment that fosters a productive dialogue. It discusses how [fact-checking](#) works and provides [resources](#) to increase media and digital literacy so people

can decide what to read, trust, and share themselves. It also features [guidelines](#) related to physical harm or violence, harmful health misinformation, voter or census interference, and informative labels for manipulated media and content digitally created or altered that may mislead.

Additionally, the firm's governance director [discusses](#) partnering with Stanford's Deliberative Democracy Lab and the Behavioral Insights Team on a community forum to discuss how generative AI chatbots should interact with people and provide guidance and advice. Meta also [discusses](#) joining Thorn and All Tech Is Human, and industry partners to mitigate potential risks of generative AI by signing on to new generative AI principles related to child safety. It also details the firm's [approach](#) to labeling AI-generated content and manipulated media, detailing its network of almost 100 independent fact-checkers who review false and misleading AI-generated content. It also [provides](#) information on labeling AI-generated images on Facebook, Instagram, and Threads.

Meta's 2023 [Human Rights Report](#) addresses AI in the context of human rights, stating that it is committed to developing and deploying AI responsibly while mitigating potential adverse human rights impacts. The firm also affirms that AI is included in its [Corporate Human Rights Policy](#), which recognizes the importance of the [OECD Principles on Artificial Intelligence](#). The report examines taking an open approach, deploying AI safely, and addressing potentially harmful generative AI outputs, including outputs that may generate potentially hateful, offensive, or discriminatory content; reinforce biases; present inaccurate information; and/or raise privacy considerations (pp.9-11).

Meta also provides a [Community Standards Enforcement Report](#) each quarter to track its progress and demonstrate its continued commitment to making its social media platforms safe and inclusive. The report discusses recent trends concerning the types of content that it prohibits and provides data for specific content policy areas. It specifically provides information on removing [fake accounts](#) and Meta's detection technology to block millions of attempts to create fake accounts every day.

Regarding oversight of this issue, the firm states in its latest proxy statement that the board, both directly and through its committees, is actively involved in overseeing risks associated with AI, including misinformation and disinformation (2024 DEF 14A, p.83). Further, the [audit and risk oversight committee](#) assists the board in overseeing the risk management of Meta and oversees certain of the firm's major risk exposures, provided that the board may, in its discretion, exercise direct oversight with respect to any such matters. The committee will review with management, at least annually, Meta's cybersecurity risk exposures and the steps management has taken to monitor or mitigate such exposures. It also reviews with management, at least annually, the firm's major ESG risk exposures and the steps management has taken to monitor or mitigate such exposures, in coordination with the other committees of the board as appropriate.

Summary

Peer Comparison

Overall, we find the Company and its peers to provide relatively commensurate disclosure with respect to their efforts to combat misinformation and disinformation generated via artificial intelligence ("AI"). All three companies maintain AI principles and board-level oversight of AI. However, only Meta maintains explicit board-level oversight of AI risks associated with misinformation and disinformation.

Analyst Note

After a similar proposal was submitted at last year's AGM, the Company published its inaugural Responsible AI Transparency Report. Additionally, it publishes a broad set of reports on its efforts to address misinformation and disinformation as a signatory to both the European Code of Practice on Disinformation and the Australian Code of Practice for Disinformation and Misinformation. However, we believe that there are certain gaps in the Company's reporting on this issue, particularly around the effectiveness of efforts to mitigate misinformation and disinformation. As such, we believe the Company could reasonably expand its disclosure to include a more specific discussion concerning how it is mitigating risks to its operations and finances as a result of misinformation and disinformation generated via artificial intelligence.

RECOMMENDATION

We understand the significant risks posed to both the Company and society resulting from a failure to adequately ensure the responsible use of artificial intelligence. This is a relatively novel and highly dynamic technology that could undoubtedly lead to unintended consequences, including the proliferation of misinformation and disinformation. We note that this is a growing risk for the Company, as it notes that it is continuing to invest in AI across its operations and is infusing generative AI capabilities into its consumer and commercial offerings. It notes that it expects AI and associated services to be a highly competitive and rapidly evolving technology, particularly as new competitors continue to enter the market. Accordingly, we believe that the Company should ensure that it is responsive to the changing regulatory and legal landscape governing this issue.

To be clear, we in no way view the Company as being unresponsive to this matter. We recognize the Company provides its Responsible AI Principles and disclosure of its policies and practices related to AI. Moreover, after a similar proposal was submitted at last year's AGM, the Company disclosed its inaugural Responsible AI Transparency Report, which

discusses how the Company builds generative applications responsibly, how it makes decisions about releasing generative applications, and how it supports its customers in building responsibly. We also recognize that the Company publishes a broad set of reports on its efforts to address misinformation and disinformation, such as the Microsoft Digital Defense Report and that every six months, it publicly reports on the steps it has taken to meet its obligations as a signatory to the European Code of Practice on Disinformation. These reports include detailed quantitative metrics on a range of measures related to disinformation. Additionally, the Company files similar reports under the Australian Code of Practice for Disinformation and Misinformation.

Despite this voluminous disclosure, we do not believe that there is adequate discussion concerning the effectiveness its policies aimed at curbing misinformation and disinformation. Understanding the efficacy of the Company's efforts in this regard is critical for shareholders who are trying to evaluate the severity of these risks and the robustness of the Company's management of this issue. Moreover, because the Company already provides extensive disclosures concerning its efforts to curb misinformation and disinformation within its AI technologies, we do not view this additional disclosure to be excessively onerous. Further, providing shareholders with an understanding of how well these initiatives are performing adds important context to this existing disclosure.

We acknowledge the complexity and rapidly evolving nature of this issue. Just as the technology is rapidly advancing, so are the regulatory, reputational and legal risks, and it is imperative that shareholders be given the tools necessary to fully understand the Company's risk exposure and how it is mitigating this exposure. Further, we note that this proposal is precatory in nature (meaning that the board has wide latitude in what is included in this report) and also specifically carves out an omission of any information that is proprietary or legally privileged. Given the Company's existing disclosure, we do not believe that production of the requested disclosure would be especially onerous given that the Company already collects and provides some disclosure of the information specified in this resolution.

In sum, we believe that the requested reporting would not be overly burdensome on the Company and that it would allow shareholders more insight into potential financial risks related to misinformation and disinformation disseminated or generated via generative AI. We believe that such information will better allow shareholders to fully consider and understand this evolving issue and incorporate these considerations into their investment decisions. As such, we believe support for this proposal is warranted at this time.

We recommend that shareholders vote **FOR** this proposal.

9.00: SHAREHOLDER PROPOSAL REGARDING REPORT ON RISKS OF AI DATA SOURCING

FOR

PROPOSAL REQUEST:	That the Company prepare a report assessing the risks presented by the real or potential unethical or improper usage of external data in AI training	SHAREHOLDER PROPONENT:	National Legal and Policy Center
BINDING/ADVISORY:	Precatory		
PRIOR YEAR VOTE RESULT (FOR):	N/A	REQUIRED TO APPROVE:	Majority of votes cast
RECOMMENDATIONS, CONCERNS & SUMMARY OF REASONING:			
FOR -	Additional disclosure will better allow shareholders to understand the Company's management of AI-related risks		

SASB MATERIALITY	PRIMARY SASB INDUSTRY: Software & IT Services
	FINANCIALLY MATERIAL TOPICS: - Environmental Footprint of Hardware Infrastructure - Recruiting & Managing a Global, Diverse & Skilled Workforce - Managing Systemic Risks from Technology Disruptions - Data Privacy & Freedom of Expression - Data Security - Intellectual Property Protection & Competitive Behavior

GLASS LEWIS REASONING

- Although we recognize the Company's existing and planned disclosures, we believe that support for this proposal could be warranted as a means to encourage the Company to ensure that its forthcoming disclosures are robust and provide a solid context for shareholders to allow them to assess the potential risks to the Company from its use of external data in the development of its AI technology.

PROPOSAL SUMMARY

Text of Resolution: *Resolved: Shareholders request the Company prepare a report, at reasonable cost, omitting proprietary or legally privileged information, to be published within one year of the Annual Meeting and updated annually thereafter, which assesses the risks to the Company's operations and finances, and to public welfare, presented by the real or potential unethical or improper usage of external data in the development and training of its artificial intelligence offerings; what steps the Company takes to mitigate those risks; and how it measures the effectiveness of such efforts.*

Proponent's Perspective

- The development and training of artificial intelligence ("AI") systems rely on vast amounts of data, and stakeholders are concerned that developers will draw from unethical or illegal sources, such as personal information collected online, copyrighted works, and proprietary commercial information provided by users;
- The Company is an early leader in AI, but shareholders are concerned regarding the Company's record on data ethics, including: (i) the Company's and OpenAI's use of personal information scraped from the web, (ii) OpenAI's recently appointing a director who was formerly head of the National Security Agency, (iii) pushback the Company received regarding its proposed AI "Recall" feature, (iv) the Company and OpenAI's inadvertent or deliberate access and use of proprietary information provided by users, and (v) the Company and OpenAI being sued by *The New York Times* for alleged copyright infringement;
- Americans surveyed by the Pew Research Center have expressed that data privacy and usage are among their main concerns with big tech's AI initiatives; and
- Prioritizing data ethics in the Company's AI development may help

Board's Perspective

- The Company has articulated publicly a number of commitments as to how it sources data for training generative artificial intelligence ("AI") models in ways that are consistent with global laws, and that respect privacy, safety, and content, and the Company communicates its approach through AI-model specific fact sheets called model cards, transparency notes, and blog posts;
- Beginning in 2025, the Company will enhance its reporting on AI to meet requirements of the EU AI Act;
- The Company uses a variety of data sources, including publicly available information, in a manner consistent with global copyright laws, and it does not source information that is behind a paywall, subject to a login or subscription (including content that violates the Company's policies), or is from a domain that has signaled a preference to opt-out of training AI models using a tag such as NO ARCHIVE;
- The Company has shared controls for website owners to make choices about the use of their content, and the Company respects standards like the robots.txt tag;
- As identified in the Company's model cards and other publicly available documentation, the Company sources the following

avoid harmful fiduciary and regulatory consequences.

The proponent has filed an [exempt solicitation](#) urging support for this proposal.

categories of data for training generative AI models responsibly: (i) publicly available data, (ii) acquired data, (iii) first party data, (iv) synthetic data, and (v) human feedback;

- In an August 2024 blog post, the Company communicated more information to consumers about how it uses chats with Copilot to help train the Company's generative AI models; and
- While this proposal focuses specifically on linking the Company's data sourcing practices to OpenAI, the Company partners with a wide variety of companies that complement its AI skills and approach, and OpenAI explained its own data sourcing practices and current work to further enhance those practices in a blog post published in May 2024.

THE PROPONENT

The National Legal & Policy Center

The proponent of this proposal is the National Legal & Policy Center ("NLPC"). The [NLPC describes](#) itself as a 501(c)(3) that "promotes ethics in public life through research, investigation, education, and legal action," and believes "the best way to promote ethics is to reduce the size of government." As NLPC is not an investor, it does not have AUM. Based on the disclosure provided by companies concerning the identity of proponents, during the first half of 2023, the NLPC submitted 24 shareholder proposals that received an average of 10.4% support, with none receiving majority shareholder support.

As part of the [corporate integrity project](#) on its website, the NLPC shares its concerns regarding "woke" corporate executives, for instance posting articles about inclusive content "[devaluing](#)" the Pixar franchise or about how the NLPC has [reported](#) Visa's chair and CEO to the SEC for ongoing "wokeness." The project also examines a supposed pushback against ESG initiatives, featuring pieces such as one describing [corporate America's anti-racism programs as racist](#) against white people and another promoting the NLPC's efforts to [nominate](#) a fossil-fuel-supporting director candidate to the board of Exxon Mobil Corporation. The NLPC has submitted other shareholder proposals that, upon first impression, appear to be consistent with environmental and social proposals that call for information or action on enhancing companies' approaches to environmental and social factors but, upon further review, appear to be designed to inhibit companies' actions in such areas.

GLASS LEWIS ANALYSIS

Glass Lewis recommends that shareholders take a close look at proposals such as this one to determine whether the actions requested by the Company will clearly lead to the protection or enhancement of long-term shareholder value. We believe it is prudent for management to assess its potential exposure to all risks, including environmental and social concerns and regulations pertaining thereto and incorporate this information into its overall business risk profile. When there is no evidence of egregious or illegal conduct that might threaten shareholder value, Glass Lewis believes that management of social issues associated with business operations are generally best left to management and directors who can be held accountable for failure to address relevant risks on these issues when they face re-election.

In this case, the proponent requests that the Company prepare a report which assesses the risks to its operations and finances, and to public welfare, presented by the real or potential unethical or improper usage of external data in the development and training of its artificial intelligence offerings, as well as what steps the Company is taking to mitigate those risks and how it measures the effectiveness of such efforts. As part of its rationale for this proposal, the proponent states that stakeholders are concerned that developers will draw from unethical or illegal sources, such as personal information collected online, copyrighted works, and proprietary commercial information provided by users. It also adds that prioritizing data ethics in the Company's AI development may help avoid harmful fiduciary and regulatory consequences.

For more information regarding the Company's and its peers' disclosures and approaches to managing the risks posed by AI, please see our analysis of Proposal 8.

COMPANY AND PEER DISCLOSURE

Company Disclosure

In its response to this proposal, the Company explains that it [communicates](#) its approach to artificial intelligence ("AI") data sourcing practices through AI-model specific fact sheets called model cards, transparency notes (2024 DEF 14A, pp.88). It states that, as identified in its model cards and other publicly available documentation, the Company sources the following categories of data for training generative AI models responsibly:

- Publicly available data: the Company trains on select publicly available data, mostly collected from

industry-standard machine learning datasets and web crawls, using practices similar to those used by search engines, and it excludes sources with paywalls, sources that contain content that violates the Company's policies, and sources that have opted out of training using web controls that the Company published. The Company also affirms that it does not train on data from domains listed in the Office of the United States Trade Representative Notorious Markets for Counterfeiting and Piracy list;

- Acquired data: the Company trains on data that it gains access to (such as archives and metadata) through negotiated arrangements with publishers and copyright owners;
- First party data: the Company trains on data from select Company consumer (but not enterprise) services in accordance with the Company's user policies and protections and with clear notice to users. For example, Bing search queries are used to improve AI models that are used to return relevant search results;
- Synthetic data: the Company sometimes trains on synthetic datasets, and the Company creates these by prompting large language models to generate specific requested output before reviewing and filtering the results to ensure they are of sufficient quality to be part of the training dataset; and
- Human feedback: the Company includes human feedback from AI trainers, red teamers, and employees in its training process, including, for example, human feedback that reinforces a quality output to a user's prompt, improving the end user experience.

(2024 DEF 14A, pp.88)

Additionally, in August 2024, the Company published a [blog post](#) on Transparency and Control in Consumer Data Use, sharing upcoming changes in how it will use consumer data, as well as its approach to ensuring that the Company's users are always in control. First, it will soon start using consumer data from Copilot, Bing, and Microsoft Start (including interactions with advertisements) to help train the generative AI models in Copilot. It states that it is making this change as real-world consumer interactions to provide greater breadth and diversity in training data, and that this will help the Company build more inclusive, relevant products, and improve the experience for all users. Second, the Company emphasizes that it will make it simple for consumers to opt-out of their data being used for training, with clear notices displayed in Copilot, Bing, and Microsoft Start. It specifies that the Company will start providing these opt-out controls in October 2024, and it won't begin training its AI models on this data until at least 15 days after the Company notifies consumers that the opt-out controls are available. It clarifies that these changes will only apply to consumers who are signed into their Microsoft Account, and asserts that consumers will retain existing controls over their data, and there are no changes to the existing options to manage how their data is used to personalize Company services for them. Further, it explains that this change only applies to how the Company will start training its generative AI models in Copilot, and it does not change existing uses of consumer data as outlined in the [Microsoft Privacy Statement](#). It continues to explain that it will gradually roll this out in different markets to ensure the Company gets this right for consumers and to comply with privacy laws around the world. For example, it will not offer this setting or conduct training on consumer data from the European Economic Area until further notice. Finally, the Company will continue to adhere to the following data protection commitments as it begins training:

- Consumer data will not be used to identify consumers;
- Consumer data will be kept private; and
- Data from minors will not be used for training.

The Company also states in its response to this proposal that it explained to consumers that the chat logs remain private when using Copilot and any information that may identify an individual will be removed before using it to help train the Company's generative AI models (2024 DEF 14A, pp.88). The Company's webpage on [Copilot](#) features a series of frequently asked questions about reviewing data use and control options. Specific questions it addresses include:

- What is a generative AI model and how is it trained?
- How does the Company use public data and other data for AI training?
- Does the Company use my data to train generative AI models?
- What data is not included in training?
- How does the Company protect my data when training generative AI models?
- What of my data is the Company using for AI training and why?
- What does the AI training opt-out mean for you?
- Who will see the AI training opt-out control in settings?
- Will my Copilot conversations become visible to other users?
- Does this change apply to Copilot Pro, Microsoft 365, or Microsoft Copilot with commercial data protection?
- Are Copilot conversations human reviewed? If yes, why?
- Can I opt-out of having my Copilot conversations human reviewed?
- Can I opt-out of AI training and still have a personalized Copilot experience?
- Does the Company share my data with third parties for AI training?

The Company also [provides](#) an [update](#) on its AI safety policies that outlines its approach to advancing responsible AI,

including by implementing voluntary commitments that it and others made at the White House convening in July 2023, and discusses its collaboration with OpenAI in this area. In the update, the Company asserts that maintaining AI products that adhere to its responsible AI commitments throughout their product lifecycle means that they are subject to an iterative cycle of mapping, measuring, and managing risk pre and post deployment. It explains that this means that policies and practices across multiple areas of inquiry on which the Company has provided context in the update, including model evaluation and red teaming, security controls, and data input controls, must be implemented iteratively as appropriate (p.20). It then discusses focusing on Data Input Controls and Audit, noting that it announced its new Copilot Copyright Commitment, which allows customers to use the Company's new Copilot services and the output they generate without worrying about copyright claims. It adds that this commitment reflects the Company's work to incorporate filters and other technologies that are designed to reduce the likelihood that Copilots return infringing content. The Company also explains that if a third party sues a commercial customer for copyright infringement for using one of the Company's Copilots or the output they generate, then the Company will defend the customer and pay the amount of any resulting adverse judgments or settlements as long as the customer has used the guardrails and content filters that the Company has built into its products (pp.22-23).

The Company [emphasizes](#) that it is committed to implementing and supporting responsible data policies and practices for inputs and outputs of AI models and applications. It states that its Responsible AI Standard and accompanying privacy, security, and accessibility standards, which apply to all AI systems that the Company develops and deploys, establish numerous data requirements impacting all of its Responsible AI principles. As a result, product teams may be required to assess the quantity and suitability of data sets, inclusiveness of data sets, representation of intended uses in training and test data, limitations to generalizability of models given training and testing data, and how they meet data collection and processing requirements, among other mandates. The Company specifies that its impact assessment and other responsible AI tools help teams conduct these assessments and provide documentation for review (p.22).

Through transparency mechanisms, it [provides](#) context to customers and other stakeholders on data processed by AI systems like the Azure OpenAI Service, including user prompts and generated content, augmented data included with prompts (i.e., for grounding), and user-provided training and validation data. It notes that customer data is also processed to analyze prompts, completions, and images for harmful content or patterns of use that may violate the Company's code of conduct or other applicable product terms. Further, the Company states that established policies dictate that training data and fine-tuned models are available exclusively for use by the customer, are stored within the same region as the Azure OpenAI resource, can be double encrypted at rest (by default with the Company's AES-256 encryption and optionally with a customer-managed key), and can be deleted by the customer at any time. Additionally, the Company clarifies that all applicable data processed by AI products is also subject to the Company's General Data Protection Regulation and other legal commitments and data privacy and security compliance offerings (p.22).

Further, the Company [discusses](#) supporting several existing technical mechanisms that content providers can use to restrict access to their data, including putting content behind a paywall or using other means to technically restrict access. It adds that content providers may also implement mechanisms to indicate that they do not intend the content that they make publicly accessible to be scraped by using machine-readable means, such as the robots.txt web standard. It also points to guardrails in its Copilots that help respect authors' copyrights. It explains that the Company has incorporated filters and other technologies that are designed to reduce the likelihood that Copilots return infringing content (pp.22-23).

Additionally, the Company provides its most recent [Responsible AI Standard](#), which discusses data governance and management. Among the data requirements for all AI systems, the Company includes:

- Defining and documenting data requirements with respect to the system's intended uses, stakeholders, and the geographic areas where the system will be deployed, and then documenting these requirements in the impact assessment;
- Defining and documenting procedures for the collection and processing of data, to include annotation, labeling, cleaning, enrichment, and aggregation, where relevant;
- If planning to use existing data sets to train the system, then assessing the quantity and suitability of available data sets that will be needed by the system in relation to the data requirements defined in the first requirement, and then documenting this assessment in the impact assessment;
- Defining and documenting methods for evaluating data to be used by the system against the data requirements defined in the first requirement; and
- Evaluating all data sets using the methods defined in the fourth requirement, and documenting the results of the evaluation.

(p.7)

The Company also discusses cyber influence operations and other cybersecurity and cybercrime issues in its [Digital Defense Report](#), and recommends limits on tools and techniques regarding the covert use of AI and theft/abuse of social media data (p.93). Additionally, the Company reviews legislative initiatives related to data-sourcing, such as the

legislation proposed in Brazil and Costa Rica that would impose on all AI systems certain security requirements, including parameters for separating and organizing training data. It also notes that China has adopted the most stringent approach to imposing security requirements on all covered AI systems, including the use of accurate and lawful training data (p.102). In discussing the seven areas of efficiencies in the Company's security operations, the Company addresses knowledge gathering from diverse external sources and states that augmenting proprietary in-house datasets with online content (such as threat intelligence and information on recent vulnerabilities) enables an organization to make better decisions. It adds that AI can scrape online content and extract security-related information at scale (p.97).

The Company disclosed its inaugural [Responsible AI Transparency Report](#) in 2024, and discusses advancements made by the Company's researchers in the emerging field of synthetic data for training and evaluating generative AI models (p.34). It also addresses building safe and responsible frontier models through partnerships and stakeholder input, explaining that it supports creators by actively engaging in consultations with sector-specific groups to obtain feedback on its tools and to incorporate their feedback into product improvements. It provides the example that news publishers expressed hesitation around their content being used to train generative AI models, but they did not want any exclusion from training datasets to affect how their content appeared in search results. The Company notes that in response to that feedback, it [launched](#) granular controls to [allow](#) web publishers to exercise greater control over how content from their websites is accessed and used (pp.30-31).

Moreover, the Company emphasizes that beginning in 2025 its commitment to transparency will expand to include the publication of a detailed summary about the content used for training its general-purpose AI models because Article 53 of the EU AI Act requires all developers of general-purpose AI models to publish such a report, and this report will be based on the upcoming template and guidance from the European Union AI Office. It explains that one stated purpose of this report is to provide evidence of how the Company is complying with the reservation of rights granted to copyright holders under the European Union's copyright laws. It also states that a specific article of the European Union copyright law permits copyright owners to exclude their works from what is referred to as Text and Data Mining activity, and using data to train AI systems can be considered a form of Text and Data Mining. The Company adds that the upcoming publication of this report reflects another step it is taking to provide details about how the Company sources data for AI in a manner that meets legal norms and promotes accountability (2024 DEF 14A, p.89).

Separate from the Company's policies, OpenAI discusses its own approach to data and AI in a May 2024 [blog post](#) on its website. For more information regarding OpenAI's approach to managing AI-related risks, please see our analysis in Proposal 8.

Regarding oversight of this issue, the [environmental, social, and public policy committee](#) assists the board in overseeing the key non-financial regulatory risks that may have a material impact on the Company and especially its ability to sustain trust with customers, employees, and the public, which includes policies and programs and related risks that concern environmental sustainability, the social and public policy impacts of technology including privacy, digital safety, and responsible artificial intelligence, and legal, regulatory, and compliance matters relating to competition / antitrust, trade, and national security. The committee also reviews and provides guidance to the board and management about key environmental and social matters such as responsible artificial intelligence, responsible sourcing, digital safety, privacy, and human rights, among others.

Peer Analysis

To compare, **Alphabet Inc.** (NASDAQ: GOOGL) states in its [privacy policy](#) that Google uses information to improve its services and to develop new products, features, and technologies that benefit its users and the public. As such, it explains that it uses publicly available information to help train Google's AI models and build products and features like Google Translate, Gemini Apps, and Cloud AI capabilities.

As part of its [AI principles](#), the firm discusses incorporating privacy design principles into the development and use of its AI technologies. It states that it will give opportunity for notice and consent, encourage architectures with privacy safeguards, and provide appropriate transparency and control over the use of data. In a [blog post](#) about building guardrails for generative AI, Alphabet discloses that it has put strong [indemnification protections](#) on both training data [used](#) for generative AI models and the generated output for users of key Google Workspace and Google Cloud services. It clarifies that if customers are challenged on copyright grounds, the firm will assume responsibility for the potential legal risks involved. It also discusses protecting against unfair bias; red-teaming to help identify current and emergent risks, behaviors, and policy violations, enabling its teams to mitigate them; implementing generative AI [prohibited use](#) policies, and safeguarding teens. Regarding protecting consumers' information, the firm states that many of the privacy protections it has had in place for years apply to its generative AI tools too and, just like other types of activity data in a Google Account, Alphabet makes it easy to pause, save, or delete it at any time, including for Bard or Search. It further adds that if a consumer chooses to use the Workspace extensions in Bard, the user's content from Gmail, Docs, and Drive is not seen by human reviewers, used by Bard to show the consumer ads, or used to train the Bard model.

Regarding board oversight, the [audit and compliance committee](#) has responsibility for oversight of risks and exposures

associated with data privacy and security, competition, legal, regulatory, compliance, civil and human rights, sustainability, and reputational risks. Alphabet also states that the board's oversight function of major risks and risk exposures, including those relating to or resulting from the firm's development and implementation of AI in its products and services, sits at the top of its risk management framework, and as set forth in Alphabet's Corporate Governance Guidelines, the board is ultimately responsible for covering strategic, financial, and execution risks and exposures associated with the firm's business strategy, production innovation, and policy and significant regulatory matters that may present material risk to its financial performance, operations, plans, prospects, or reputation. Further, it adds that the board's skills and expertise, including deep technical expertise in computer science, facilitates oversight of a highly complex global business, and the full board meetings have regularly and extensively covered AI issues. The audit committee and senior management provide the board with reports and updates regarding issues and risk exposures regarding AI development, and these discussions ensure that the board is fully involved in the oversight of Alphabet's business strategies and plans as they relate to AI (2024 DEF 14A, p.94).

To further compare, **Meta Platforms, Inc. (NASDAQ: META)** [provides](#) its commitment to responsible AI, with details on its five pillars, which include that protecting the privacy and security of people's data is the responsibility of everyone at the firm, and that people who use Meta's products should have more transparency and control around how data about them is collected and used. The firm's [privacy policy](#) addresses how it uses consumers' information, specifying that it uses information it has, information from researchers and datasets from publicly available sources, professional groups, and non-profit groups to conduct and support research. The firm [links](#) to additional information about its efforts to research and innovate for social good, explaining that some examples of its research include, among other things, supporting research in areas like artificial intelligence and machine learning. It then provides resources for consumers to access and manage their information.

Meta specifically [addresses](#) how it uses information for generative AI models and features. It states that since it takes such a large amount of data to teach effective models, a combination of sources are used for training. It explains that Meta uses information that is publicly available online and licensed information, and it also uses information shared on Meta's Products and services. It notes that this information could be things like posts or photos and their captions, but clarifies that it does not use the content of consumers' private messages with friends and family to train its AIs. The firm states that when it collects public information from the internet or licenses data from other providers to train its models, it may include personal information. For example, if Meta collects a public blog post it may include the author's name and contact information, but it adds that when it does get personal information as part of this public and licensed data that it uses to train its models, it doesn't specifically link this data to any Meta account. Further, it explains that even if someone doesn't use Meta's products and services or have an account, the firm may still process information about that person to develop and improve AI at Meta. This could happen if a person appears anywhere in an image shared on the firm's products or services by someone who does use them or if a consumer mentions information about that person in posts or captions that they share on Meta's products and services.

Additionally, the firm [asserts](#) that it is committed to being transparent about the legal bases that it uses for processing information, and that it believes use of this information is in the legitimate interests of Meta, its users, and other people. In the European region and the UK, the firm relies on the basis of legitimate interests to collect and process any personal information included in the publicly available and licensed sources to develop and improve AI at Meta. For other jurisdictions where applicable, it relies on an adequate legal basis to collect and process this data. The firm also [provides](#) additional information regarding data subject rights for third party information used for AI at Meta.

Regarding specific products, Meta provides a [Responsible Use Guide](#) for Meta Llama, and its [Meta Llama Guard](#) approach to trust and safety, also detailing its Prompt Guard to protect LLM powered applications from things like prompt injections, or inputs that exploit the inclusion of untrusted data from third parties into the context window of a model to get it to execute unintended instructions. The firm also [discusses](#) essential data and additional data it collects for its Smart Glasses.

Regarding oversight of this issue, the firm states that the board, both directly and through its committees, is actively involved in overseeing risks associated with AI, including misinformation and disinformation (2024 DEF 14A, p.83). Further, the [audit and risk oversight committee](#) assists the board in overseeing the risk management of Meta and oversees certain of the firm's major risk exposures, provided that the board may, in its discretion, exercise direct oversight with respect to any such matters. The committee will review with management, at least annually, Meta's cybersecurity risk exposures and the steps management has taken to monitor or mitigate such exposures. It also reviews with management, at least annually, the firm's major ESG risk exposures and the steps management has taken to monitor or mitigate exposures, in coordination with the other committees of the board as appropriate.

Summary

Peer Comparison

Overall, we find the Company and its peers to provide meaningful disclosure of their usage of external data in the development and training of their AI offerings, and their measures to mitigate associated risks.

Analyst Note

In response to this proposal, the Company states that beginning in 2025 it will publish a detailed summary about the content used for training its general-purpose AI models to provide evidence of how the Company is complying with the reservation of rights granted to copyright holders under the EU's copyright laws. The report will provide details about how the Company sources data for AI in a manner that meets legal norms and promotes accountability. Further, the board reviews and provides guidance to management about key environmental and social matters such as responsible AI, responsible sourcing, digital safety, privacy, and human rights.

RECOMMENDATION

As discussed in Proposal 8, the Company's investment in and use of AI technologies presents a significant growth opportunity for the Company. However, alongside this opportunity also comes significant risks. The Company faces a variety of legal, regulatory and reputational risks on account of this issue, some of which stem from the usage of external data in the development and training of its AI offerings, the topic of this resolution.

We acknowledge that the Company provides some information concerning this issue. More importantly, it also notes in response to this proposal that, beginning in 2025 it will publish a detailed summary about the content used for training its general-purpose AI models to provide evidence of how the Company is complying with the reservation of rights granted to copyright holders under the EU's copyright laws. It further states that the report will provide details about how the Company sources data for AI in a manner that meets legal norms and promotes accountability. We believe that this could go a long way to addressing the request of this resolution.

Although we recognize the Company's existing and planned disclosures, we still believe that support for this proposal could be warranted as a means to encourage the Company to ensure that its forthcoming disclosures are robust and provide a solid context for shareholders to allow them to assess the potential risks to the Company from its use of external data in the development of its AI technology. Accordingly, we believe that shareholders should vote in favor of this proposal at this time.

We recommend that shareholders vote **FOR** this proposal.

COMPETITORS / PEER COMPARISON

	MICROSOFT CORPORATION	ALPHABET INC.	APPLE INC.	AMAZON.COM, INC.
Company Data (MCD)				
Ticker	MSFT	GOOGL	AAPL	AMZN
Closing Price	\$415.00	\$172.49	\$225.00	\$202.61
Shares Outstanding (mm)	7,434.9	12,241.0	15,115.8	10,515.0
Market Capitalization (mm)	\$3,085,475.5	\$2,119,197.7	\$3,401,060.2	\$2,130,446.4
Enterprise Value (mm)	\$3,161,473.5	\$2,128,527.7	\$3,490,176.2	\$2,213,890.4
Latest Filing (Fiscal Period End Date)	09/30/24	09/30/24	09/28/24	09/30/24
Financial Strength (LTM)				
Current Ratio	1.3x	1.9x	0.9x	1.1x
Debt-Equity Ratio	0.34x	0.09x	2.09x	0.61x
Profitability & Margin Analysis (LTM)				
Revenue (mm)	\$254,190.0	\$339,859.0	\$391,035.0	\$620,128.0
Gross Profit Margin	69.3%	58.1%	46.2%	48.4%
Operating Income Margin	44.5%	32.1%	31.5%	9.8%
Net Income Margin	35.6%	27.7%	24.0%	8.0%
Return on Equity	35.6%	32.1%	157.4%	22.6%
Return on Assets	14.6%	16.5%	21.5%	7.1%
Valuation Multiples (LTM)				
Price/Earnings Ratio	34.3x	22.9x	37.0x	43.3x
Total Enterprise Value/Revenue	12.4x	6.3x	8.9x	3.6x
Total Enterprise Value/EBIT	28.0x	19.5x	28.3x	36.5x
Growth Rate* (LTM)				
5 Year Revenue Growth Rate	14.4%	17.0%	8.5%	18.5%
5 Year EPS Growth Rate	18.0%	26.5%	15.4%	32.9%
Stock Performance (MCD)				
1 Year Stock Performance	12.2%	27.5%	18.6%	39.6%
3 Year Stock Performance	21.6%	15.1%	42.5%	9.6%
5 Year Stock Performance	176.0%	161.4%	237.0%	131.2%

Source: Capital IQ

MCD (Market Close Date): Calculations are based on the period ending on the market close date, 11/18/24.

LTM (Last Twelve Months): Calculations are based on the twelve-month period ending with the Latest Filing.

*Growth rates are calculated based on a compound annual growth rate method.

A dash ("-") indicates a datapoint is either not available or not meaningful.

VOTE RESULTS FROM LAST ANNUAL MEETING DECEMBER 7, 2023

Source: 8-K (sec.gov) dated December 8, 2023

RESULTS

NO.	PROPOSAL	FOR	AGAINST/WITHHELD	ABSTAIN	GLC REC
1.1	Elect Reid G. Hoffman	99.03%	0.69%	0.28%	For
1.2	Elect Hugh F. Johnston	91.13%	8.69%	0.18%	Against
1.3	Elect Teri L. List	97.83%	2.00%	0.17%	For
1.4	Elect Catherine MacGregor	99.60%	0.23%	0.17%	For
1.5	Elect Mark Mason	99.61%	0.21%	0.18%	For
1.6	Elect Satya Nadella	93.96%	5.61%	0.43%	For
1.7	Elect Sandra E. Peterson	97.97%	1.86%	0.17%	For
1.8	Elect Penny S. Pritzker	99.44%	0.38%	0.17%	For
1.9	Elect Carlos A. Rodriguez	97.16%	2.67%	0.17%	For
1.10	Elect Charles W. Scharf	98.31%	1.52%	0.18%	For
1.11	Elect John W. Stanton	99.33%	0.50%	0.18%	For
1.12	Elect Emma N. Walmsley	98.86%	0.97%	0.17%	For
2.0	Advisory Vote on Executive Compensation	93.34%	6.19%	0.47%	For
4.0	Ratification of Auditor	95.11%	4.71%	0.18%	For

FREQUENCY OF ADVISORY VOTE ON EXECUTIVE COMPENSATION

NO.	PROPOSAL	1 YEAR	2 YEARS	3 YEARS	ABSTAIN	GLC REC
3.0	Frequency of Advisory Vote on Executive Compensation	98.44%	0.13%	1.25%	0.19%	1 Year

SHAREHOLDER PROPOSALS*

NO.	PROPOSAL	FOR	AGAINST	GLC REC
5.0	Shareholder Proposal Regarding Report on Median Compensation and Benefits Related to Reproductive and Gender Dysphoria Care	1.01%	98.99%	Against
6.0	Shareholder Proposal Regarding EEO Policy Risk Report	0.82%	99.18%	Against
7.0	Shareholder Proposal Regarding Report on Government Takedown Requests	1.78%	98.22%	Against
8.0	Shareholder Proposal Regarding Risks of Developing Military Weapons	15.20%	84.80%	For
9.0	Shareholder Proposal Regarding Report on Climate Risk In Employee Retirement Options	8.89%	91.11%	Against
10.0	Shareholder Proposal Regarding Report on Tax Transparency	21.26%	78.74%	For
11.0	Shareholder Proposal Regarding Report on Siting in Countries of Significant Human Rights Concern	33.58%	66.42%	Against
12.0	Shareholder Proposal Regarding Third-Party Political Expenditures Reporting	5.37%	94.63%	Against
13.0	Shareholder Proposal Regarding Report on AI Misinformation and Disinformation	21.17%	78.83%	Against

*Abstentions excluded from shareholder proposal calculations.

APPENDIX

GLASS LEWIS PEERS VS PEERS DISCLOSED BY COMPANY

GLASS LEWIS	MSFT
Alphabet Inc.*	Accenture PLC
Apple Inc.*	Adobe Inc.
Amazon.com, Inc.*	Cisco Systems, Inc.
Meta Platforms, Inc.*	Intel Corporation
Johnson & Johnson*	Merck & Co., Inc.
Verizon Communications Inc.*	Oracle Corporation
AT&T Inc.*	Pfizer Inc.
The Walt Disney Company*	The Procter & Gamble Company
International Business Machines Corporation*	QUALCOMM Incorporated
Walmart Inc.	Salesforce, Inc.
Comcast Corporation*	Netflix, Inc.
The Boeing Company	
UnitedHealth Group Incorporated	
NVIDIA Corporation*	
Tesla, Inc.*	
*ALSO DISCLOSED BY MSFT	

QUESTIONS

Questions or comments about this report, GL policies, methodologies or data? Contact your client service representative or go to www.glasslewis.com/public-company-overview/ for information and contact directions.

DISCLAIMERS

© 2024 Glass, Lewis & Co., and/or its affiliates. All Rights Reserved.

This Proxy Paper report is intended to provide research, data and analysis of proxy voting issues and, therefore, is not and should not be relied upon as investment advice. Glass Lewis analyzes the issues presented for shareholder vote and makes recommendations as to how institutional shareholders should vote their proxies, without commenting on the investment merits of the securities issued by the subject companies. Therefore, none of Glass Lewis' proxy vote recommendations should be construed as a recommendation to invest in, purchase, or sell any securities or other property. Moreover, Glass Lewis' proxy vote recommendations are solely statements of opinion, and not statements of fact, on matters that are, by their nature, judgmental. Glass Lewis research, analyses and recommendations are made as of a certain point in time and may be revised based on additional information or for any other reason at any time.

The information contained in this Proxy Paper report is based on publicly available information. While Glass Lewis exercises reasonable care to ensure that all information included in this Proxy Paper report is accurate and is obtained from sources believed to be reliable, no representations or warranties express or implied, are made as to the accuracy or completeness of any information included herein. Such information may differ from public disclosures made by the subject company. In addition, third-party content attributed to another source, including, but not limited to, content provided by a vendor or partner with whom Glass Lewis has a business relationship, as well as any [Report Feedback Statement](#) or Partner Insights attached to this Proxy Paper report, are the statements of those parties and shall not be attributed to Glass Lewis. Neither Glass Lewis nor any of its affiliates or third-party content providers shall be liable for any losses or damages arising from or in connection with the information contained herein, or the use of, or inability to use, any such information.

This Proxy Paper report is intended to serve as a complementary source of information and analysis for subscribers in making their own voting decisions and therefore should not be relied on by subscribers as the sole determinant in making voting decisions. Glass Lewis expects its subscribers to possess sufficient experience and knowledge to make their own decisions entirely independent of any information contained in this Proxy Paper report. Subscribers are ultimately and solely responsible for making their own voting decisions, including, but not limited to, ensuring that such decisions comply with all agreements, codes, duties, laws, ordinances, regulations, and other obligations applicable to such subscriber.

All information contained in this Proxy Paper report is protected by law, including, but not limited to, copyright law, and none of such information may be copied or otherwise reproduced, repackaged, further transmitted, transferred, disseminated, redistributed or resold, or stored for subsequent use for any such purpose, in whole or in part, in any form or manner or by any means whatsoever, by any person without Glass Lewis' express prior written consent.

This report should be read and understood in the context of other information Glass Lewis makes available concerning, among other things, its research philosophy, approach, [methodologies](#), sources of information, and [conflict management, avoidance and disclosure policies and procedures](#), which information is incorporated herein by reference. Glass Lewis recommends all clients and any other consumer of this Proxy Paper report carefully and periodically evaluate such information, which is available at: <http://www.glasslewis.com>.

PARTNER INSIGHTS

The pages following this appendix are included with this Proxy Paper report for informational purposes only. They contain data and insights produced by Glass Lewis' strategic business partners and none of the information included therein is a factor in Glass Lewis' analyses or vote recommendations.

About ESG Book

ESG Book is a global leader in sustainability data and technology. Launched in 2018, the company offers a wide range of sustainability-related data, scoring, and technology products that are used by many of the world's leading investors and companies. Covering over 35,000 companies, ESG Book's product offering includes ESG raw data, company-level and portfolio-level scores and ratings, analytics tools, and a SaaS data management and disclosure platform. ESG Book's solutions cover the full spectrum of sustainable investing including ESG, climate, net-zero, regulatory, and impact products. Read more on: www.esgbook.com

SUSTAINALYTICS ESG PROFILE

ESG Risk Rating

Negligible **Low** Med High Severe

All data and ratings provided by:

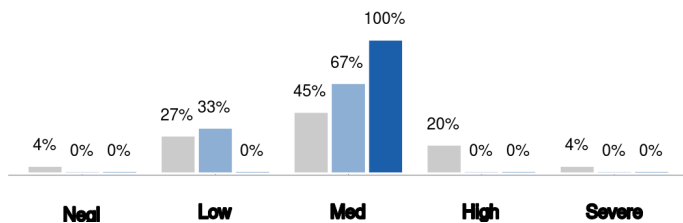


Data Received On: **November 06, 2024**

Rating Overview

The company is at low risk of experiencing material financial impacts from ESG factors, due to its low exposure and strong management of material ESG issues. The company is noted for its strong corporate governance performance, which is reducing its overall risk. The company is noted for its strong stakeholder governance performance, which is reducing its overall risk. However, the company has a high level of controversies.

ESG Risk Rating Distribution



Relative Performance

	Rank*	Percentile*
Global Universe	1451 of 15035	11th
Software & Services (Industry Group)	59 of 958	7th
Enterprise and Infrastructure Software (Subindustry)	16 of 378	5th

* 1st = lowest risk

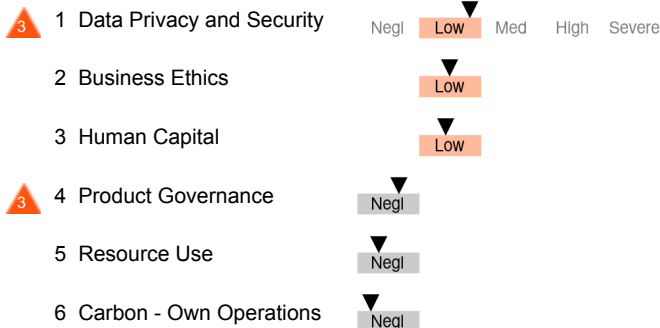
Exposure to ESG Risk



Management of ESG Risk



Top Material Issues



▲ = Noteworthy Controversy Level



Risk Details

Exposure		The company's sensitivity or vulnerability to ESG risks.
Company Exposure		
Management		
Manageable Risk		Material ESG risk that can be influenced and managed through suitable policies, programmes and initiatives.
Managed Risk		Material ESG risk that has been managed by a company through suitable policies, programmes or initiatives.
Management Gap		Measures the difference between material ESG risk that could be managed by the company and what the company is managing.
Unmanageable Risk		Material ESG risk inherent in the products or services of a company and/or the nature of a company's business, which cannot be managed by the company.
ESG Risk Rating		
Overall Unmanaged Risk		Material ESG risk that has not been managed by a company, and includes two types of risk: unmanageable risk, as well as risks that could be managed by a company through suitable initiatives but which may not yet be managed.

NOTEWORTHY CONTROVERSIES

SEVERE

The Event has a severe impact on the environment and society, posing serious business risks to the company. This category represents exceptional egregious corporate behavior, high frequency of recurrence of incidents, very poor management of ESG risks, and a demonstrated lack of willingness by the company to address such risks.

- No severe controversies

HIGH

The Event has a high impact on the environment and society, posing high business risks to the company. This rating level represents systemic and/or structural problems within the company, weak management systems and company response, and a recurrence of incidents.

- No high controversies

SIGNIFICANT

The Event has a significant impact on the environment and society, posing significant business risks to the company. This rating level represents evidence of structural problems in the company due to recurrence of incidents and inadequate implementation of management systems or the lack of.

- Data Privacy and Security

NO PRODUCT INVOLVEMENT



Alcoholic Beverages



Oil Sands



Arctic Drilling



Genetically Modified Plants & Seeds



Pesticides



Adult Entertainment



Gambling



Tobacco



Controversial Weapons



Thermal Coal

* Range values represent the percentage of the Company's revenue. N/A is shown where Sustainalytics captures only whether or not the Company is involved in the product.

DISCLAIMER

Copyright © 2024 Sustainalytics. All rights reserved.

Sustainalytics' environmental, social and governance ("ESG") data points and information contained in the ESG profile or reflected herein are proprietary of Sustainalytics and/or its third parties suppliers (Third Party Data), intended for internal, non-commercial use, and may not be copied, distributed or used in any way, including via citation, unless otherwise explicitly agreed in writing. They are provided for informational purposes only and (1) do not constitute investment advice; (2) cannot be interpreted as an offer or indication to buy or sell securities, to select a project or make any kind of business transactions; (3) do not represent an assessment of the issuer's economic performance, financial obligations nor of its creditworthiness.

These are based on information made available by third parties, subject to continuous change and therefore are not warranted as to their merchantability, completeness, accuracy or fitness for a particular purpose. The information and data are provided "as is" and reflect Sustainalytics' opinion at the date of their elaboration and publication. Sustainalytics nor any of its third-party suppliers accept any liability for damage arising from the use of the information, data or opinions contained herein, in any manner whatsoever, except where explicitly required by law. Any reference to third party names or Third Party Data is for appropriate acknowledgement of their ownership and does not constitute a sponsorship or endorsement by such owner. A list of our third-party data providers and their respective terms of use is available on our website.

For more information, visit <http://www.sustainalytics.com/legal-disclaimers>.

This ESG profile is presented for informational purposes and is not a factor in Glass Lewis' analyses or vote recommendations.

All data and ratings provided by:



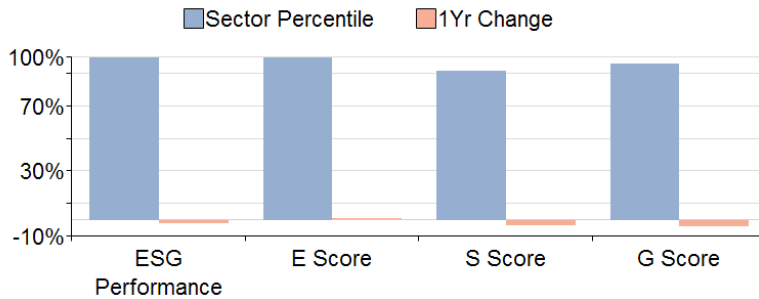
SUSTAINALYTICS

<https://www.sustainalytics.com/>

ESG BOOK PROFILE

Summary of ESG Performance Score

All data and ratings provided by:



esgbook

www.esgbook.com

Country:	United States
Sector:	Technology Services
Industry:	Packaged Software
Data Received:	3/7/2024

ESG Performance Score Details

The ESG Performance Score provides investors and corporates with a systematic and comprehensive sustainability assessment of corporate entities. The score measures company performance relative to salient sustainability issues across the spectrum of environmental, social and governance. The score is driven by a sector-specific scoring model that emphasises financially material issues, where the definition of financial materiality is inspired by the Sustainability Accounting Standards Board (SASB). For more detail please see the [ESG Performance Score methodology here](#).

ESG Performance Score		Environmental	Social	Governance
Absolute Score	76.1	Score 92.6	68.7	68.5
Sector Percentile	100.0%	Weight	31.1%	41.7%
1 Year Change	-2.1%	Sector Percentile	100.0%	91.6%
2 Year Change	8.3%	1 Year Change	0.1%	-3.4%
3 Year Change	7.5%			-3.8%

Risk Score Details

The Risk Score provided by ESG Book assesses company exposures relative to universal principles of corporate conduct defined by the UN's Global Compact. The score is accompanied by a transparent methodology and full data disclosure, enabling users to comprehend performance drivers, explain score changes, and explore associated raw data. Tailored for both investors and corporates, it serves as a universe selection tool for investors identifying companies more exposed to critical sustainability issues, while corporates can use it to assess their exposures, conduct peer comparisons, and pinpoint disclosure gaps. For more detail please see the [risk score methodology user guide here](#).

Risk Score		Human Rights	Labour Rights	Environment	Anti-corruption
Absolute Score	80.2	Score 75.2	91.0	88.0	66.6
Sector Percentile	100.0%	Weight	25.0%	25.0%	25.0%
1 Year Change	-1.7%	Sector Percentile	99.8%	98.8%	99.7%
2 Year Change	8.1%	1 Year Change	-1.3%	-2.5%	-1.4%
3 Year Change	13.3%			-1.4%	-1.5%

Business Involvements - Over a 5% Revenue Threshold

ESG Book has not found any business involvements for the Company that exceed a 5% revenue threshold.

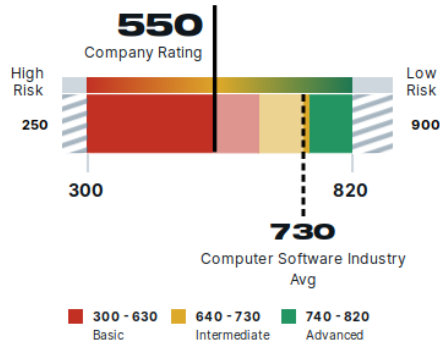
© ESG Book GmbH 2024 (together with its branch and subsidiary companies, "ESG Book") is a limited liability company organized under the laws of Germany, with registered number HRB 113087 in the commercial register of the court of Frankfurt am Main, and having its seat and head office at Zeppelinallee 15, 60325 Frankfurt am Main, Germany. All rights reserved. The "ESG Book Profile" is provided "as is" and does not constitute investment advice or a solicitation or an offer to buy any security or instrument or to participate in investment services. ESG Book makes no representation or warranty, express or implied, as to the accuracy or completeness of the information contained herein, and accepts no liability for any loss, of whatever kind, howsoever arising, in relation thereto. ESG Book shall not be responsible for any reliance or decisions made based on information contained within the ESG Book Profile. This ESG Book Profile is presented for informational purposes and is not a factor in Glass Lewis' analyses or vote recommendations.

BITSIGHT CYBERSECURITY RATING PROFILE

Microsoft Group of Companies

COMPARATIVE INDUSTRY:
Computer Software

Bitsight Security Rating



Likelihood of Ransomware Incidents

6.4x as Likely vs a 750+ company



Source: [Link to Research](#)

Likelihood of Data Breach Incidents

3x as Likely vs a 700+ company



Source: [Link to Research](#)

What is a BitSight Security Rating?

BitSight Security Ratings are a measurement of a company's security performance over time. BitSight Security Ratings are generated through the analysis of externally observable data, leveraging BitSight's proprietary techniques to identify the scope of a company's entire digital footprint. BitSight continuously measures security performance based on evidence of compromised systems, diligence, user behavior, and data breaches to provide an objective, evidence-based measure of performance. This data-driven approach requires no cooperation from the rated company. The Rating is representative of the cybersecurity performance of an entire company, including its subsidiaries, business units, and geographic locations. ¹ **This company has delegated security**

EXECUTIVE REPORT

All data and ratings provided by:

Data Received on: **Nov 19, 2024**

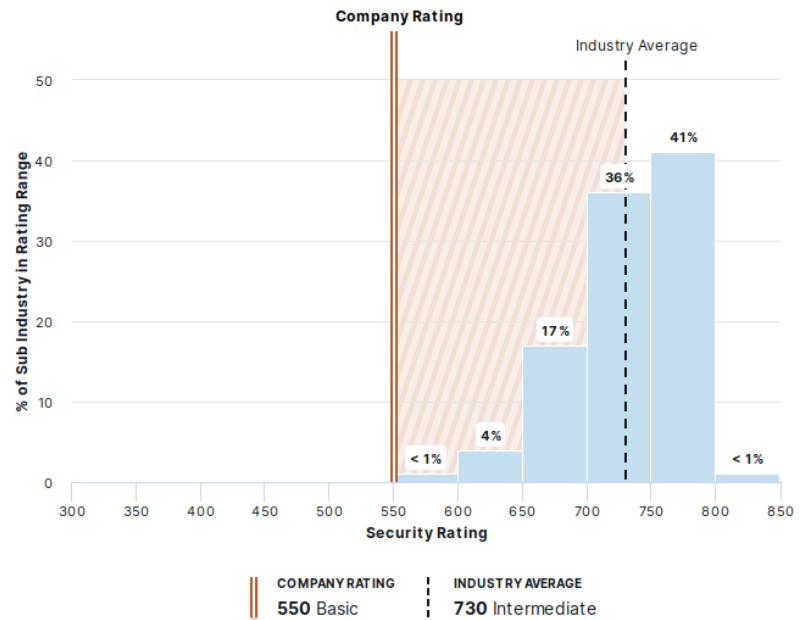
BITSIGHT

PEER ANALYTICS

This compares a company against its industry:

TOTAL COMPANIES
35,185

INDUSTRY RATING
Bottom 1% of the industry

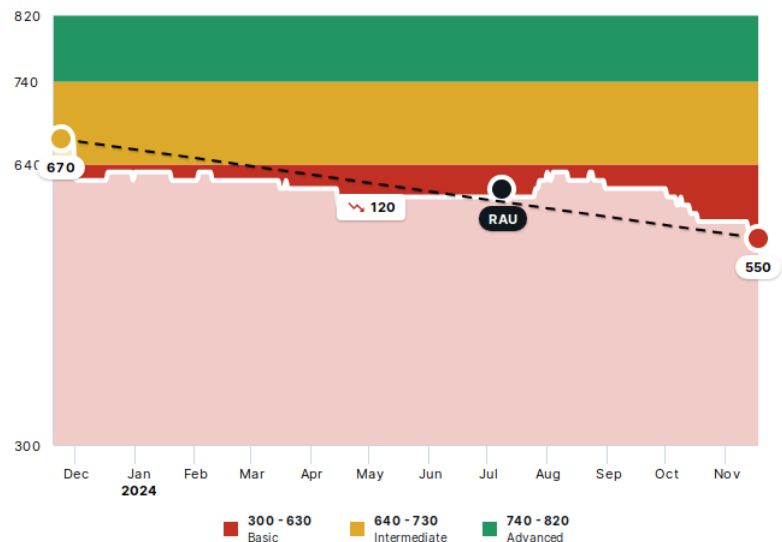


PERFORMANCE OVER THE LAST 12 MONTHS

This rating change graph includes all rating changes events, including but not limited to, publicly disclosed security events.

HIGHEST
670 on Nov 24, 2023

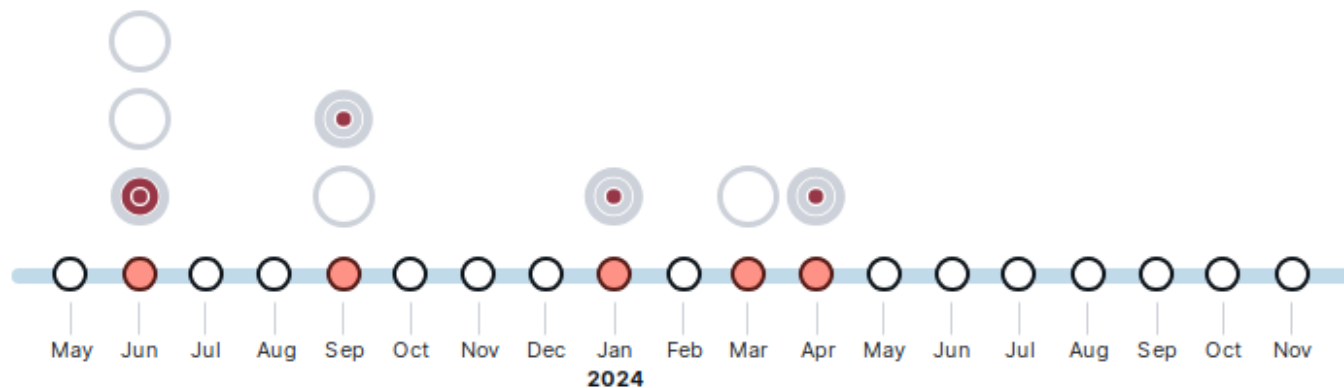
LOWEST
550 on Nov 18, 2024



controls, see disclaimer below for more details.

PUBLICLY DISCLOSED SECURITY INCIDENTS THE LAST 18 MONTHS

Security incidents are publicly disclosed events of unauthorized access, often involving data loss or theft. These events are graded based on several factors, including the number of data records lost or exposed.



General Security Incident

Severity: Minor ? ? ? ?

Public Discovery on Apr 9, 2024

Microsoft Azure—a subsidiary of Microsoft Group of Companies: Microsoft Azure left confidential information accessible online.

General Security Incident

Severity: Informational ? ? ? ?

Public Discovery on Mar 25, 2024

GitHub, Inc.—a subsidiary of Microsoft Group of Companies: An unauthorized party gained access to a Top.gg GitHub user account. This action compromised the information of an unknown number of individuals.

General Security Incident

Severity: Minor ? ? ? ?

Public Discovery on Jan 19, 2024

Several Microsoft employee email accounts were subject to unauthorized access. This action compromised the company's source code repositories and internal systems.

General Security Incident

Severity: Minor ? ? ? ?

Public Discovery on Sep 18, 2023

Due to an error, Microsoft Corporation leaked private data while publishing a bucket of open-source AI training data on GitHub.

General Security Incident

Severity: Informational ? ? ? ?

Public Discovery on Sep 6, 2023

A Microsoft employee account was subject to unauthorized access. This action compromised confidential information.

Other Disclosure

Severity: Informational ? ? ? ?

Public Discovery on Jun 9, 2023

Microsoft OneDrive—a subsidiary of Microsoft Group of Companies: Microsoft OneDrive suffered a denial-of-service attack that compromised the availability of some users accounts.

Other Disclosure

Severity: Informational ? ? ? ?

Public Discovery on Jun 6, 2023

Microsoft Outlook—a subsidiary of Microsoft Group of Companies: Microsoft Outlook fell victim to a DDOS attack that compromised the availability of the systems.

General Security Incident

Severity: Moderate ? ? ? ?

Public Discovery on Jun 4, 2023

Nuance Communications, Inc.—a subsidiary of Microsoft Group of Companies: Due to a vulnerability in Progress Software's MOVEit Cloud, the personal information of an unknown number of individuals was accessible online. This includes data held by Progress for its customers who use the SaaS service.

ADDITIONAL INFORMATION

Security Rating Overview

BitSight Security Ratings are a measurement of a company's security performance over time. BitSight Security Ratings are generated through the analysis of externally observable data, leveraging BitSight's proprietary techniques to identify the scope of a company's entire digital footprint. BitSight continuously measures security performance based on evidence of compromised systems, diligence, user behavior, and data breaches to provide an objective, evidence-based measure of performance. This data-driven approach requires no cooperation from the rated company. The Rating is representative of the cybersecurity performance of an entire company, including its subsidiaries, business units, and geographic locations.

In some cases, a company may designate one or more subsidiaries, business units or locations as representative of the company's overall digital footprint. In these cases, BitSight flags those companies in its reports as a Primary Rating, meaning that the company has undertaken this optional step in further articulating its digital footprint.

Companies often use Primary Ratings to exclude parts of their digital infrastructure that may not be useful in describing their cyber risk and resulting security posture. As examples, Primary Ratings often exclude guest wireless networks, security test environments, or networks used for customer hosting. BitSight does not validate Primary Ratings or whether the digital assets organizations exclude in creating Primary Ratings are properly excluded, nor does it validate the predictive quality of Primary Ratings. Go to [this web page](#) for more information about Primary Ratings.

BitSight rates companies on a scale of 250 to 900, with 250 being the lowest measure of security performance and 900 being the highest. A portion of the upper and lower edge of this range is currently reserved for future use. The effective range as of this report's generation is 300-820. Go to [this web page](#) to learn more about how BitSight security ratings are calculated.

1 Delegated Security Controls Companies ("DCEs")

Companies that, in the normal course of business, delegate the control and allow their customers or the public to control some aspects of security for some of their IT assets can be classified as companies with delegated security controls. These companies may provide internet access or network resources-as-a-service, conduct internet scanning or threat research, or perform other business activities that place some aspects of cyber security beyond the companies' control.

To more accurately assess the security posture of companies with delegated security controls, Bitsight excludes findings in assets with delegated controls from the companies' security ratings. The same enterprise Bitsight ratings algorithm with the same set of risk vectors are used. This means that if you monitor or do business with them, you can consider their Bitsight rating to be a better measurement of their external security performance as a result of this updated approach.

Rating Algorithm Update (RAU)

BitSight periodically makes improvements to its ratings algorithm. These updates often include new observation capabilities, enhancements to reflect the rapidly changing threat landscape, and adjustments to further increase quality and correlation with business outcomes. BitSight's Rating and Methodology Governance Board governs these changes so that they adhere to BitSight's principles and policies. BitSight also has a Policy Review Board which reviews and arbitrates customer disputes associated with its ratings. More information about the Policy Review Board and its cases can be found [here](#). Additionally, BitSight provides a preview of ratings algorithm changes customers (and what the likely impact will be) well before they affect the live ratings, inviting comments and feedback on these changes.

Publicly Disclosed Security Incidents

The Security Incidents risk vector involves a broad range of events related to the unauthorized access of a company's data. BitSight collects information from a large number of verifiable sources such as news organizations and regulatory reports obtained via Freedom of Information Act requests or local analogs. This risk vector only impacts BitSight Security Ratings if a confirmed incident occurs. For more information about publicly disclosed security incidents and how BitSight ratings are calculated, [please go here](#).

Disclaimer

© 2024 Bitsight Technologies, Inc. (together with its majority owned subsidiaries, "Bitsight"). All rights reserved. This report and all the data contained herein (the "Information") is the proprietary information of Bitsight. Information is provided on an "as is" basis, for an organization's internal use and informational purposes only, and does not constitute investment or financial advice, nor recommendations to purchase, sell, or hold particular securities. Bitsight hereby disclaims any and all warranties whatsoever, including, but not limited to, any warranties of merchantability or fitness for a particular purpose with respect to the Information. Bitsight shall not be responsible for any reliance or decisions made based upon Information, and to the extent permitted by law, shall not be liable for any direct, indirect, incidental, consequential, special, or punitive damages associated therewith. Except as otherwise permitted in an applicable underlying agreement, this report may not be reproduced in whole or in part by any means of reproduction.