

RECENT TRENDS IN GENERATIVE ARTIFICIAL INTELLIGENCE LITIGATION IN THE UNITED STATES

Date: 5 September 2023

US Litigation and Dispute Resolution Alert

By: Christopher J. Valente, Michael J. Stortz, Amy Wong, Peter E. Soskin, Michael W. Meredith

Although still in their infancy, a growing number of recently-filed lawsuits associated with generative artificial intelligence (AI) training practices, products, and services have provided a meaningful first look into how US courts may address the privacy, consumer safety, and intellectual property protection concerns that have been raised by this new, and inherently evolving, technology. The legal theories that have served as the basis of recent claims have varied widely, are often overlapping, and have included invasion of privacy and property rights; patent, trademark, and copyright infringement; libel and defamation; and violations of state consumer protection laws, among others.

To date, courts have appeared reluctant to impose liability on AI developers and have expressed skepticism of plaintiffs' rhetoric around AI's purported world-ending potential. Courts have also found a number of recent complaints to be lacking in the specific, factual, and technical details necessary to proceed beyond the pleadings stage. This alert aims to provide a snapshot of the current litigation landscape in the rapidly growing field of generative-AI law in the United States.

CURRENT LANDSCAPE

Privacy Cases

Over a two-week period in June and July, 2023, a number of federal class action lawsuits were filed in the US District Court for the Northern District of California, many by the same law firm, against the developers of some of the most well-known generative AI products on the market, including OpenAI, Inc. (OpenAI) and Alphabet Inc./Google LLC (Google).

On 28 June 2023, for example, in *P.M. v. OpenAI LP*,¹ an anonymous group of plaintiffs filed suit against OpenAI LP (OpenAI) and Microsoft, Inc. (Microsoft) alleging that OpenAI stole private and personal information belonging to millions of people by collecting publicly-available data from the Internet² to develop and train its generative AI tools—including ChatGPT (an AI text generator), Dall-E (an AI image generator), and Vall-E (an AI speech generator). The plaintiffs allege that OpenAI's practice of using datasets of information gathered from the Internet to train its generative AI tools constitutes theft, misappropriation, and a violation of their privacy and property rights. The complaint also includes claims for violations of the Electronic Communications Privacy Act, the Computer Fraud and Abuse Act (CFAA), various state consumer protection statutes, and a number of common law claims. Emphasizing arguments that AI poses a danger to human civilization in the form of misinformation, malware, and even autonomous weapons systems, the plaintiffs seek injunctive relief, including the implementation of human oversight and human-developed ethical protocols for OpenAI's products. In addition,

plaintiffs seek various forms of class-wide damages and/or restitution on behalf of multiple classes that purport to include all persons whose personal information was used without permission.

On 11 July 2023, in *J.L. v. Alphabet Inc.*,³ the same plaintiffs' firm as in *P.M.* filed a similar class action complaint against Google, asserting both privacy and copyright law violations. As with *P.M.*, the plaintiffs raise theoretical concerns relating to the proliferation of AI.⁴ The plaintiffs in *J.L.* argued that many of Google's generative AI products, including Bard (a text generator), Imagen and Gemini (two text-to-image diffusion models), MusicLM (a text-to-music tool), and Duet AI (a data visualization tool), all relied on training data that Google collected from the Internet. The complaint does note that Google was transparent and disclosed its data gathering practices, but suggested that Google should have explored other options for the development of training data, such as purchasing from the commercial data market. Additionally, and without any specific evidence, the plaintiffs argued that Google violated the Copyright Act because: (1) Google's AI products allegedly used copyrighted text, music, images, and data for training purposes; and (2) the AI products themselves, as well as their expressive output, constitute infringing derivative works. Like the plaintiffs in *P.M.*, the plaintiffs in *J.L.* are seeking broad injunctive relief aimed at restricting Google's generative AI products. These plaintiffs are also seeking further specific relief with respect to their claims for copyright infringement, as well as various forms of class-wide damages related to its theories.

Copyright Cases

Numerous cases have been filed against generative AI developers asserting violations of copyright law. Foundationally, these claims amount to a challenge to developers' use of data collected from the Internet to train generative AI models—and whether collecting and using publicly-available data that may be subject to copyright protection constitutes infringement. To date, these cases have been aimed at generative AI developers that have engaged in the data collection and generative AI training process. The viability of these claims remains to be seen because the use of copyrighted materials only for training purposes does not involve the impermissible “copying” or “reproduction” for commercial purposes that are traditionally contemplated by copyright law and, as more fully described below, may fall squarely within the definition of a “fair use” of the information, which is expressly permitted by the Copyright Act.

One of the most well-known cases alleging copyright infringement is *Andersen v. Stability AI Ltd.*⁵ In that case, plaintiffs Sarah Andersen, Kelly McKernan, and Karla Ortiz, on behalf of a putative class of artists, alleged that Stability AI, Ltd. and Stability AI, Inc. (collectively, Stability AI) and others scraped billions of copyrighted images from online sources, without permission, in order to train their image-generating models to produce seemingly new images without attribution to the original artists who supplied the training material. They further argued that this practice deprived artists of commissions and allowed the defendants to profit from the artists' copyrighted works. In their motion to dismiss, the defendants argued that the models do not copy or store any images, copyrighted or otherwise. Rather, the defendants explained, their models only analyze the properties of online images to generate parameters that were later used to assist the model in creating new and unique images from text prompts, as opposed to reproducing or copying any portion of the underlying images used for training.

At a hearing on the defendants' motion to dismiss on 19 July 2023, Judge William Orrick expressed skepticism regarding the plaintiffs' claims indicating he would tentatively dismiss them. Specifically, he explained that: (1) the images produced by the models are not “substantially similar” to plaintiffs' art; and (2) because the models had been trained on “five billion compressed images” it is “implausible that [plaintiffs'] works are involved” in the

creation of those images. Judge Orrick did, however, provide plaintiffs with an opportunity to amend their complaint “to provide more facts” proving otherwise.⁶

GitHub, Inc. (GitHub), the well-known online code repository, is also the subject of a putative class action filed in November 2022 in the Northern District of California under the caption *Doe v. GitHub, Inc.*⁷ In that case, the anonymous plaintiffs are developers who allegedly published licensed code on GitHub's website and claim that GitHub used that code to train its AI-powered coding assistant, Copilot. The developer-plaintiffs sued GitHub, Microsoft, and OpenAI alleging violations of privacy and property rights, including violation of copyright management laws based on GitHub's purported use of licensed materials without appropriate attribution.

In the motion to dismiss, defendants argued, among other things, that plaintiffs could not plausibly allege that any code that they individually authored was impermissibly used by Copilot because that model functions by generating its own unique coding suggestions based on what it learned from reviewing open source code without copying or reproducing any portion of that open source code. GitHub also addressed head-on plaintiffs' allegation (based on an internal study) that “about 1% of the time” Copilot generated snippets of code similar to the publicly available code from which it learned. Even if this were true, GitHub argued, the plaintiffs could not allege that their own code fell within that 1%. Said differently, the plaintiffs could not connect the dots between their own code and any generated code. Therefore, GitHub explained, none of the plaintiffs in *Doe* could ever have the requisite standing to pursue any legal claim, copyright-based or otherwise. As detailed below, this may be a viable defense available to generative AI developers in nearly all of the legal challenges they presently face relating to data collection and generative AI model training.

In its decision on the motion to dismiss, the court dismissed the plaintiffs' privacy rights claims because, as alleged, the plaintiffs had not met their burden to demonstrate injury-in-fact sufficient to confer standing. With respect to their property right claims, the court determined that the plaintiffs failed to establish an injury-in-fact sufficient to confer standing for their claims for damages; however, found that their allegations could plausibly give rise to standing to pursue injunctive relief because the plaintiffs adequately alleged a danger that their code could potentially be output by Copilot in violation of the license. The court went on to address certain of plaintiffs' individual claims under a Rule 12(b)(6) standard. The court granted plaintiffs leave to re-plead all of but two of their claims and plaintiffs filed an amended complaint in June 2023. GitHub subsequently filed a motion to dismiss the amended complaint as well. The court has not yet ruled on that motion.

It is interesting to note that the plaintiffs in *Doe* did not assert any direct claims for copyright infringement, instead relying on a theory of improper copyright information management; namely, that defendants violated the Digital Millennium Copyright Act by failing to provide the appropriate attribution, copyright notice, or license terms when making use of the plaintiffs' code. In contrast, on 28 June 2023, two authors, Paul Tremblay and Mona Awad, filed a class action lawsuit against OpenAI in *Tremblay v. OpenAI, Inc.*, in which they directly assert copyright infringement on behalf of a class of authors.⁸ Their infringement claims were premised on the theory that: (1) the authors' copyrighted material was duplicated and ingested as part of OpenAI's training data; and (2) that OpenAI's large language models and their output constitute infringing derivative works.⁹ The plaintiffs further argued that the models' outputs, specifically summaries of the authors' copyrighted books, are likewise derivative and infringing works. The plaintiffs claim that their lawsuit is not intended to limit the functioning of OpenAI's algorithms, generally speaking, and is instead aimed only at compensating authors for the value that their books added to the AI's training data. As such, they are seeking actual and statutory damages, together with injunctive

relief. Two weeks later, authors Sarah Silverman, Christopher Golden, and Richard Kadrey filed similar class action suits against OpenAI and others in *Silverman v. OpenAI, Inc.*¹⁰ and *Kadrey v. Meta Platforms, Inc.*¹¹

Notably, copyright claims arising out of the functioning of AI products are not new. One of the earliest lawsuits asserting copyright infringement by an AI product was *Thomson Reuters Enterprise Centre GmbH v. ROSS Intelligence Inc.*,¹² which was filed in the US District Court for the District of Delaware on 6 May 2020. In that case, plaintiff Thomson Reuters, owner of the Westlaw legal research platform, argued that ROSS Intelligence, Inc. (ROSS) copied Westlaw's legal database without permission, in order to train its competing AI-powered legal research software. On 26 April 2022, Thomson Reuters' claims survived a motion to dismiss. Cross motions for summary judgment were filed in December 2022 that went to the heart of ROSS' primary defense in the matter—that its training protocols constitute fair use because only unprotected ideas and legal decisions contained in Thomson Reuters' database were used for training, as opposed to Westlaw's indexing and searching systems, which have been granted copyright protection. Those motions remain pending.

More copyright cases may soon follow by other content creators, including newspapers and other media organizations, many of which are currently in the process of negotiating a licensing arrangement permitting the use of their reporting, articles, and others publications by generative AI developers. Should those negotiations break down, litigation will likely follow. *The New York Times*, for example, is reportedly considering legal action against OpenAI because ChatGPT has, in the past, incorporated *The New York Times*' articles and reporting when responding to users' inquiries into the news and other topics, in essence “becoming a direct competitor with the paper” and greatly reducing the need for users to “visit the publisher's website” at all, which might result in economic losses for the paper.¹³ Additionally, on 21 August 2023, it was reported that *The New York Times* had blocked OpenAI's web crawler, thus preventing OpenAI from using *The New York Times*' content to train AI models.¹⁴

Trademark Cases

In February 2023, in *Getty Images (US), Inc. v. Stability AI, Inc.*,¹⁵ the media company, Getty Images, Inc. (Getty), filed suit against Stability AI asserting claims of copyright and trademark infringement. Specifically, the complaint asserts that Stability AI “scraped” Getty's website for images and data used in the training of its image-generating model, Stable Diffusion, with the aim of establishing a competing product or service. Getty's complaint includes a number of unique intellectual property claims. Like other plaintiffs, Getty has alleged that: (1) Stability AI reproduced Getty's copyrighted material in connection with the training of its Stable Diffusion model; and (2) the model creates infringing derivative works as output. However, because Stable Diffusion has generated images that include a modified version of Getty's watermark, Getty's complaint also includes allegations that: (1) the inclusion of Getty's watermark constitutes trademark infringement likely to result in confusion as to owner of the images produced, as well as trademark dilution due to the low-quality images often produced by Stability AI's model; and (2) that by removing its watermarks for training or applying a version of its watermarks to Stable Diffusion's output, Stability AI has provided false copyright information in violation of 17 U.S.C. § 1202(a). Getty has also argued that these actions constitute unfair competition in violation of Delaware's Uniform Deceptive Trade Practices Act, and seeks damages together with an order that Stability AI destroy any Stable Diffusion models trained using Getty's content. Stability AI has moved to dismiss Getty's complaint on jurisdictional and substantive grounds or to transfer the case to the US District Court for the Northern District of California. Stability AI's motion remains pending.

Right of Publicity and Facial Recognition Cases

In contrast to some of the headline-grabbing lawsuits broadly challenging the legality or safety of AI technology on the whole, a few narrow cases involving the use of AI in connection with facial recognition software may end up being better received by courts. The reason is, in part, that these cases can more clearly identify specific individuals that have suffered harm—namely, those that have had their faces scanned and analyzed without permission.

For example, in April 2023, in *Young v. NeoCortex, Inc.*,¹⁶ television personality Kyland Young filed a class action complaint in the US District Court for the Central District of California against software developer NeoCortex, Inc. (NeoCortex) claiming that NeoCortex's AI-powered “Reface” application, which allows users to digitally “swap” their faces with celebrities and public figures in photos and videos, constitutes a violation of the common law right of publicity, protected by California's Right of Publicity Statute. In response, NeoCortex sought to dismiss the complaint asserting that plaintiffs' state law claims are barred by the First Amendment and pre-empted by the federal Copyright Act. This case remains pending.

Similar claims were asserted in *Flora v. Prisma Labs, Inc.*,¹⁷ a February 2023 suit filed in the US District Court for the Northern District of California, in which a putative class of Internet users alleged that Prisma Labs, Inc.'s portrait-generating application, Lensa, scanned the facial information of Internet users without their consent, in violation of Illinois' data privacy statute.

Tort Cases

In June 2023, radio host and public figure Mark Walters filed the first tort case against an AI company.¹⁸ Walters sued OpenAI for libel after its ChatGPT generated a fabricated complaint containing allegations against Walters for fraud and embezzlement. The suit stems from a journalist that asked ChatGPT to provide a summary of a pending civil rights lawsuit filed in the US District Court for the Western District of Washington.¹⁹ In its response, ChatGPT indicated that Walters was a defendant in that suit and stands accused of fraud and embezzlement, which was untrue. Walters' suit claims that ChatGPT's false summary constitutes libel inappropriately published to ChatGPT's users. In its motion to dismiss, OpenAI argued that, because AI-generated content is probability-based, it is widely known that ChatGPT occasionally provides false information, a phenomenon known as “hallucinations”, and that the platform itself is incapable of actual malice, as is required to establish a claim for libel against a public figure like Walters. Finally, OpenAI argued that merely providing a response to a prompt does not constitute a “publication” within the meaning of libel law. That motion remains pending, and questions remain regarding whether Section 230 of the Communications Decency Act, which shields Internet firms from liability for information produced by a third party and hosted on their platforms, applies in this case.

POSSIBLE DEFENSES

There is a variety of defenses that have already been effectively asserted by defendants in generative-AI litigation. Common themes include lack of standing, reliance on the “fair use” doctrine, and the legality of so-called “data scraping.” The following is a brief summary of the key principles underlying each of these possible defenses that AI developers may rely on in future litigation.

Lack of Standing

In July 2023, the Seventh Circuit Court of Appeals in *Dinerstein v. Google LLC*²⁰ affirmed the dismissal of breach of privacy claims brought on behalf of a putative class of patients of the University of Chicago Medical Center (UMMC) for lack of standing. The *Dinerstein* decision could provide AI developers with precedent for an important, and possibly complete, defense to claims that rely on the assumption that mere use of copyrights, consumer, or personal data to “train” AI models constitutes a legally cognizable harm. The holding in *Dinerstein* suggests, to the contrary, and regardless of the legal theory selected, the individual owners of copyrighted, personal, or private data used in AI “training” must demonstrate a plausible, concrete injury to establish standing to pursue those theories.

In *Dinerstein*, the plaintiffs alleged that UMMC breached its contractual privacy arrangements with its patients, invaded their privacy, and violated Illinois' consumer protection statute by using several years of anonymized patient medical records to train an AI model that could be used in software that could be used to anticipate patients' future healthcare needs. The US District Court for the Northern District of Illinois ultimately dismissed the plaintiff's claims due to lack of standing and for failure to state a claim, noting that plaintiffs failed to establish damages associated with the disclosure of their anonymized patient data or defendants' tortious intent.

Affirming the dismissal, the Seventh Circuit Court of Appeals “agreed with [the] decision to dismiss the case” but indicated that the analysis should “begin[] and end[] with standing.”²¹ Specifically, it explained that, because plaintiffs failed to allege that any patient data was used to identify any specific member of the class and the Defendants contractually and “explicitly agreed not to identify any individual,”²² plaintiffs could not establish the existence of any “concrete and particularized, actual or imminent” harm necessary to “supply the basis for standing.”²³

This rationale underlies the arguments found in the motions to dismiss (and the court's decisions thereon) in *Andersen* and *Doe* with respect to standing; namely, that if plaintiffs cannot demonstrate that their specific information or work product was improperly used by the model in question, their claims must fail for failure to establish injury-in-fact.

“Fair Use”

Another broad defense that might be successfully pursued by AI developers against any copyright claim is the well-recognized doctrine of “fair use.” Fair use is a defense to claims of infringement when copyrighted material is used in a “transformative” way. Transformative use can occur when copyrighted material is used to serve different market “functions” or expand the “utility” of the copyrighted work. The doctrine appears particularly appropriate for the AI training process, which does not involve the traditionally impermissible copying and commercial reproduction of copyrighted work and, instead, only analyzes copyrighted material to detect patterns in an effort to develop a new “function” or “application”, namely, a large language model or other generative AI product.

To date, no US court has explained the appropriate application of the “fair use” doctrine in the context of generative AI models or AI-generated materials. However, the doctrine has provided a complete defense in similar situations. For example, in *Authors Guild v. Google, Inc.*,²⁴ the Second Circuit Court of Appeals concluded that a search engine's publication of small portions of copyrighted books was transformative because it improved access to that information. The Ninth Circuit, in *Kelly v. Arriba Soft Corp.*,²⁵ held the same with respect to searchable images of copyrighted visual artwork.

In response to lawsuits alleging copyright infringement, some AI developers have already suggested the fair use doctrine's applicability. In its copyright dispute with Thomson Reuters, for example, ROSS expressly asserted this defense, arguing that it was using Thomson Reuters' legal database for an entirely different purpose—to train its model to write new and unique code. Some have suggested that the Supreme Court's decision in *Google v. Oracle*,²⁶ which determined that Google's use of portions of Oracle's code to create its Android operating system was fair use, may provide support for ROSS' theory.

GitHub and Microsoft have also argued that the plaintiffs in *Doe v. GitHub, Inc.* affirmatively chose not to assert claims of copyright infringement because they “would run headlong into the doctrine of fair use.”²⁷ Stability AI, similarly, has also defended their model's training processes by stating, “anyone that believes that this isn't fair use does not understand the technology and misunderstands the law.”²⁸

Legality of So-Called “Data Scraping”

Finally, while generative AI developers may have relied on scraping of the Internet to develop training datasets for their products, they are far from the first group of companies to “scrape” the Internet for commercially useful information. In fact, it is a common practice amongst data science and technology companies. One such company, hiQ Labs, Inc., for example, famously “scraped” information from the publicly available profiles of online users of the business-networking site LinkedIn in order to provide employers with data and analysis regarding potential recruits and their job-seeking behaviors. In *hiQ Labs, Inc. v. LinkedIn Corp.*,²⁹ the Ninth Circuit Court of Appeals rejected claims that the practice of “scraping” publicly available data constitutes an invasion of privacy or violation of the CFAA. In its decision, the court focused on the distinction of publicly available data and data marked “private,” and held that accessing publicly available data does not constitute access without authorization under CFAA unless the data has been marked “private.”

AI developers will likely be able to take advantage of the precedent established in *hiQ Labs* to defend their data collection practices and can further expect that the *hiQ Labs* decision will likely feature prominently in the forthcoming motions to dismiss in the *P.M.* and *J.L.* cases pending in the US District Court for the Northern District of California.

FUTURE TRAJECTORY

While the outcomes of these early generative AI cases are far from certain, preliminary indications suggest that courts are not succumbing to the hype and rhetoric, and are approaching generative AI claims with a healthy level of skepticism. Yet, many of the potential defenses have still not been tested in the context of generative AI. The coming months will be pivotal in setting the tone for generative AI litigation moving forward.

In addition, while the current wave of generative AI litigation continues to work its way through the courts, recent trends suggest that plaintiffs' attorneys may be eager to expand beyond the generative AI developers to target companies that adopt or use generative AI products or solutions. As such, businesses exploring or using generative AI products and services would do well to ensure they understand the technology they are adopting as well as the risks associated with how that technology works.

FOOTNOTES

¹ Case No. 3:23-cv-3199 (N.D. Cal. 2023).

² Plaintiffs refer to the process of “scraping” data from the Internet. Scraping is a commonly used, automated process, employed by data science and technology companies, to collect large amounts of information from publicly available online sources. This data is then used to compile datasets, which are in turn used to “train” AI models.

³ Case No. 3:23-cv-03440 (N.D. Cal. 2023).

⁴ The same law firm represents the plaintiffs in both P.M. and J.L.

⁵ Case No. 3:23-cv-00201 (N.D. Cal. 2023).

⁶ Minute Entry, Andersen, Case No. 23-cv-00201 (N.D. Cal. 2023); Blake Brittain, *US Judge Finds Flaws in Artists' Lawsuit Against AI Companies*, *REUTERS* (July 19, 2023), <https://www.reuters.com/legal/litigation/us-judge-finds-flaws-artists-lawsuit-against-ai-companies-2023-07-19/>.

⁷ Case No. 4:22-cv-06823 (N.D. Cal. 2022).

⁸ Case No. 4:23-cv-03223 (N.D. Cal. 2023).

⁹ On 28 August 2023, OpenAI filed a motion to dismiss this claim, and others, in both the Silverman and Tremblay matters, arguing that the plaintiffs' derivative infringement claims are based on the “erroneous legal conclusion that every single ChatGPT output is necessarily an infringing 'derivative work' – which is a very specific term in copyright law – because those outputs are, in only a remote and colloquial sense, 'based on' an enormous training dataset that allegedly included plaintiffs' books.” The motion does not seek to dismiss the plaintiffs' claim for 'direct' copyright infringement, which OpenAI has indicated it will “seek to resolve as a matter of law at a later stage of the case”.

¹⁰ Case No. 3:23-cv-03416 (N.D. Cal. 2023).

¹¹ Case No. 3:23-cv-03417 (N.D. Cal. 2023).

¹² Case No. 1: 20-cv-00613 (D. Del. 2021).

¹³ Bobby Allyn, *New York Times' considers legal action against OpenAI as copyright tensions swirl*, NPR (Aug. 16, 2023), <https://www.npr.org/2023/08/16/1194202562/new-york-times-considers-legal-action-against-openai-as-copyright-tensions-swirl>.

¹⁴ Jay Peters and Wes Davis, *The New York Times blocks OpenAI's web crawler*, *THE VERGE* (August 21, 2023), <https://www.theverge.com/2023/8/21/23840705/new-york-times-openai-web-crawler-ai-gpt>

¹⁵ Case No. 1:23-cv-00135 (D. Del. 2023).

¹⁶ Case No. 2:23-cv-02496 (C.D. Cal. 2023).

¹⁷ Case No. 3:23-cv-00680 (N.D. Cal. 2023).

¹⁸ *Walters v. OpenAI*, Case No. 23-cv-03122 (N.D. Ga. 2023).

¹⁹ *The Second Amend. Found. v. Robert Ferguson.*, Case No. 2:2023-cv-00647 (W.D. Wash. 2023).

²⁰ *Dinerstein v. Google, LLC*, 73 F.4th 502, 508 (7th Cir. 2023)

²¹ *Id.*

²² *Id.* at 516.

²³ *Id.* at 508.

²⁴ 804 F.3d 202 (2d cir. 2015).

²⁵ 336 F.3d 811 (9th Cir. 2002).

²⁶ 141 S. Ct. 1163 (2021).

²⁷ Motion to Dismiss Operative Complaint in Consolidated Actions, *Doe v. GitHub, Inc.*, Case No. 4:22-cv-06823 (N.D. Cal. 2022), ECF No. 108.

²⁸ Riddhi Setty, *First AI Art Generator Lawsuits Threaten Future of Emerging Tech*, BLOOMBERG LAW (January 20, 2023), <https://news.bloomberglaw.com/ip-law/first-ai-art-generator-lawsuits-threaten-future-of-emerging-tech>.

²⁹ Case No. 17-3301 (9th Cir. 2022).

KEY CONTACTS



CHRISTOPHER J. VALENTE
PARTNER

BOSTON, WASHINGTON DC
+1.617.951.9071
CHRISTOPHER.VALENTE@KLGATES.COM



AMY WONG
PARTNER

ORANGE COUNTY
+1.949.623.3599
AMY.WONG@KLGATES.COM



MICHAEL W. MEREDITH
ASSOCIATE

SEATTLE
+1.206.370.7614
MICHAEL.MEREDITH@KLGATES.COM



MICHAEL J. STORTZ
PARTNER

SAN FRANCISCO
+1.415.882.8011
MICHAEL.STORTZ@KLGATES.COM



PETER E. SOSKIN
ASSOCIATE

SAN FRANCISCO
+1.415.882.8046
PETER.SOSKIN@KLGATES.COM

This publication/newsletter is for informational purposes and does not contain or convey legal advice. The information herein should not be used or relied upon in regard to any particular facts or circumstances without first consulting a lawyer. Any views expressed herein are those of the author(s) and not necessarily those of the law firm's clients.