

グローバルビジネスと人権: 「責任ある AI ガバナンス」ニュースレター 第 1 号: AI の仕組みとガバナンスの地図を描く ——技術理解からはじめる、なぜ複数の規範が併存するのか

2026 年 7 月

One Asia Lawyers Group
コンプライアンス・ニュースレター
アジア ESG/SDGs プラクティスグループ

ひとことで言うと

生成 AI は、大量の文章データから「次にどの言葉が来やすいか」を統計的に学習し、その学習結果(パラメータ)を使って入力に対する出力を確率的に作り出す仕組みである。この仕組みを知らないまま規制だけを読むと、規制が何を対象に、何を狙っているのかが見えてこない。また、AI ガバナンスを構成する規範には、法律(ハードロー)、政府指針(ソフトロー)、国際原則、認証規格など複数の種類があり、国によってその配合比率が大きく異なる。本号では、まず AI の技術的な姿を平易に押さえたうえで、これらの規範を一枚の地図として整理する。この地図は、第 2 号以降、毎号繰り返し参照する共通言語となる。

はじめに——なぜ技術理解から始めるのか

「責任ある AI ガバナンス」シリーズの第 1 号は、あえて法律の話から始めない。多くの AI ガバナンス論は規制の名称や条文番号から入り、統治の対象である技術そのものの説明を後回しにしがちである。しかし、規範の射程や妥当性を評価する作業は、「その規範が何を統治しようとしているのか」という技術的実体の理解と切り離せない。例えば「学習データ」と一言でいっても、著作物を含むか、個人情報を含むかによって適用される法令は変わる。「モデルの重みを公開する」という一つの経営判断は、検証可能性・安全性・競争政策・輸出管理という互いに異質な論点へ同時に波及する(第 4 号で詳述する)。技術の解像度が低いままでは、規制の解像度も上がらない——これが本シリーズを技術理解から始める理由である。本号は、実務判断に必要な最低限の技術的語彙を提供したうえで、AI ガバナンスの全体構造を一枚の地図として提示する。

第 1 部 生成 AI はどう動いているのか

1. 学習(トレーニング)と推論(インファレンス)

現在主流の生成 AI(大規模言語モデル)は、大きく分けて 2 つの工程で動いている。第一が学習(トレーニング)であり、インターネット上の文章や書籍など、膨大な量のテキストデータを読み込ませ、「ある単語列の次に、どのような単語が来やすいか」という統計的な規則性をモデルに反映させる工程である。第二が推論(インファレンス)であり、学習を終えたモデルが、利用者からの入力(プロンプト)を受け取り、学習した規則性に基づいて、確率的に最も自然な続きの単語を一つずつ生成していく工程である。

比喩的に言えば、学習は「大量の文章を読み込んで言葉の使い方の感覚を身につける」段階、推論は「その感覚をもとに、その場で文章を組み立てる」段階にあたる。AI が時として事実と異なる内容をもっともらしく生成してしまう現象(いわゆるハルシネーション)は、この仕組みの故障ではなく帰結である——モデルは事実を「記憶」して取り出しているのではなく、統計的にもっともらしい続きを生成しているに過ぎない。この点を押さえておくと、多くの規範が AI の出力を「無検証で信頼してよい情報源」ではなく「人間による確認を前提とした支援ツール」として位置づけている理由が、技術の側から自然に理解できる。

2. パラメータ・重みとは何か

学習の結果としてモデルの内部に蓄積される、大量の数値の集合を「パラメータ」または「重み(ウェイト)」と呼ぶ。現在の最先端モデルは、数千億から 1 兆を超える規模のパラメータを持つものもある。パラメータの数が多いほど、より複雑な言語パターンを表現できる可能性が高まるが、同時に学習・運用に必要な計算資源(コンピューティングリソース)も増大する。

重要なのは、この重みの集合自体には、人間が読んで理解できるような明示的なルールや説明が書き込まれているわけではないという点である。「なぜこの出力が生成されたのか」を、個々のパラメータの値から人間が直接読み解くことは事実上できない。これが、AI システムが「ブラックボックス」と呼ばれる所以であり、EU AI Act の透明性要求や NIST AI RMF の説明可能性に関する議論など、多くの規範が「説明可能性」「透明性」を重視する技術的な背景でもある。

3. オープンウェイトモデルとクローズドモデル

AI 開発企業が、学習済みモデルの重みそのものを一般に公開し、誰でもダウンロード・改変・自社運用できるようにする方式を「オープンウェイト」と呼ぶ。これに対し、重みを非公開とし、API(アプリケーション・プログラミング・インタフェース)等を通じてのみ利用させる方式を「クローズドモデル」と呼ぶ。

この違いは、単なる技術選択にとどまらず、統治構造上の意味を持つ。オープンウェイトモデルは、第三者が独自に安全性を検証したり、特定用途向けに改変(ファインチューニング)したりすることを可能にする一方、ひとたび公開されると提供者側が事後的に利用を制限することが事実上困難になる。クローズドモデルは、提供者が利用状況にある程度把握・制御できる一方、検証可能性は提供者の自己申告や開示姿勢に依存する。中国発のオープンウェイトモデルの急速な台頭は、この力学を世界規模で可視化させた象徴的な出来事であり、詳細は第 4 号で扱う。

用語	ひとことで言うと
学習(トレーニング)	大量のテキスト等のデータから、単語や文の出現パターンを統計的に学習させる工程。
推論(インファレンス)	学習済みのモデルが、入力を受けて実際に出力(回答)を生成する工程。
パラメータ／重み	学習によって調整される数値群。モデルの「知識」や「振る舞い」の実体に相当する。
ハルシネーション	AI が、事実と異なる内容を、もっともらしい形で生成してしまう現象。
オープンウェイト	学習済みモデルの重みを公開し、誰でもダウンロード・改変・検証できる方式。
クローズドモデル	重みを非公開とし、API 等を通じてのみ利用させる方式。
ファインチューニング	公開済みの学習済みモデルに追加の学習を行い、特定用途に最適化する工程。

第 2 部 AI ガバナンスの地図を描く

1. AI ガバナンスとは何か——リスクベース・アプローチという共通言語

AI ガバナンスとは、AI の開発・提供・利用に伴うリスクを管理しつつ、その便益を最大限に引き出すための制度・体制・実務の総体を指す。その基底にあるのが「リスクベース・アプローチ」——すべての AI を一律の強度で規制するのではなく、想定されるリスクの大き

さに応じて規制と対応の強度を変えるという発想——であり、この点では世界の主要な規範は珍しく一致している。ただし後述のとおり、その「リスク」を誰が・いつ・どう評価するかという設計になると、法域ごとの違いが一気に表面化する。共通言語は入口だけであり、中身は同床異夢である。

2. 規範の 5 類型

AI ガバナンスを構成する規範は、性質に応じて大きく 5 つに類型化できる。

1. 条約・国際原則——OECD AI 原則、G7 広島 AI プロセス国際指針・国際行動規範など、法的拘束力を持たないが国際的な合意を形成する文書群。
2. 国内法(ハードロー)——国会で制定され、法的拘束力と(多くの場合)罰則を伴う法律。EU AI Act が典型例。
3. 政府ガイドライン(ソフトロー)——法的拘束力はないが、行政が事業者の行動指針として示すもの。日本の AI 事業者ガイドラインが典型例。
4. 国際規格(認証可能な基準)——ISO/IEC 42001 のように、第三者認証の対象となりうる管理システム規格。
5. 企業独自の自主基準——個々の AI 開発企業が任意に公表する安全性ポリシー等。法的拘束力はないが、市場や投資家からの評価対象となる。

3. なぜ国によってハードローとソフトローの配合比率が異なるのか

同じ「AI ガバナンス」という言葉を使っている、各国・地域が採用する規範の組み合わせは大きく異なる。この違いは偶然ではなく、少なくとも 3 つの構造的な要因から説明できる。第一に、法文化の違いである。包括的な立法によって事前にルールを固める大陸法的な伝統を持つ法域は成文のハードローに向かいやすく、判例と個別執行を重ねる英米法的な伝統や、行政指導を通じた柔軟な調整を好む法域はソフトローや既存法の運用に向かいやすい。第二に、自国の AI 産業の位置取りである。規制の名宛人となる大手 AI 開発企業を自国内に抱える国は、イノベーション阻害への警戒から拘束的な規制に慎重になりやすく、主として域外企業の AI を「利用する側」に立つ法域は、厳格な事前規制を課すことへの国内的な抵抗が相対的に小さい。第三に、規制を域外に輸出する意思と能力の差である。EU が単一のハードローを選んだ背景には、GDPR で実証された「自域内の市場規模を梃子に自らの規制を事実上の世界標準に押し上げる」という戦略(いわゆるブリュッセル効果)の再現を狙う面があり、この戦略は巨大な単一市場と一元的な立法機構を持つ主体にしか採り得ない。つまり各法域の規範の配合比率は、その国の法文化・産業構造・国際戦略を映す鏡であり、だからこそ簡単には収斂しない。

法域	基本アプローチ	特徴
EU	ハードロー中心	EU AI Act(Regulation(EU)2024/1689)により、リスク分類に応じた義務・罰則を法律で直接定める。
米国	分裂的・州主導	連邦レベルの統一法はなく、大統領令による方針転換と、カリフォルニア州等の州法が事実上の規律を形成する。
日本	基本法+ソフトロー	AI 法(令和 7 年法律第 53 号)という基本法に、AI 事業者ガイドライン等のソフトローを組み合わせる「二層構造」モデルを採る。

日本の AI 法は、2025 年 5 月 28 日に成立、6 月 4 日に公布され、AI 基本計画・AI 戦略本部に関する規定を含め同年 9 月 1 日に全面施行された。同法は、企業に直接的な義務を課す条項を実質的に第 7 条(活用事業者の責務)に限定しており、欧州のような網羅的な義務・罰則体系とは異なる「基本法+理念法」としての性格を持つ。この構造の詳細は第 2 号で扱う。

4. 「リスクベース」という言葉の多義性

「リスクベース・アプローチ」という言葉自体は、EU・米国・日本のいずれの規範体系でも用いられているが、その内実は同じではない。EU AI Act のリスクベースは、AI システムの用途をあらかじめ「許容できないリスク」「高リスク」「限定的リスク」「最小リスク」の4区分に分類し、区分ごとに異なる法的義務を課するという、いわば静的・分類的なリスクベースである。これに対し NIST AI RMF のリスクベースは、Govern(統治)・Map(特定)・Measure(評価)・Manage(対応)という4つの機能を継続的に循環させる、動的・プロセス的なリスクマネジメントの考え方に近い。両者を同じ「リスクベース」という言葉で語ると、議論がかみ合わなくなることがあるため、この違いを意識する必要がある。

5. グローバル展開企業のための規制環境クイックマップ——米国・欧州・中国・インド・東南アジア

本ニューズレターの読者の多くは、日本国内にとどまらず、米国・欧州・中国・インド・東南アジアなど複数の法域で事業を展開する企業の関係者であろう。前節で見たとおり、各法域の規範の配合比率はその国の法文化・産業構造・国際戦略を反映して異なり、しかも2026年に入ってから目まぐるしく動いている。とはいえ、個々の法令を暗記する必要はない。

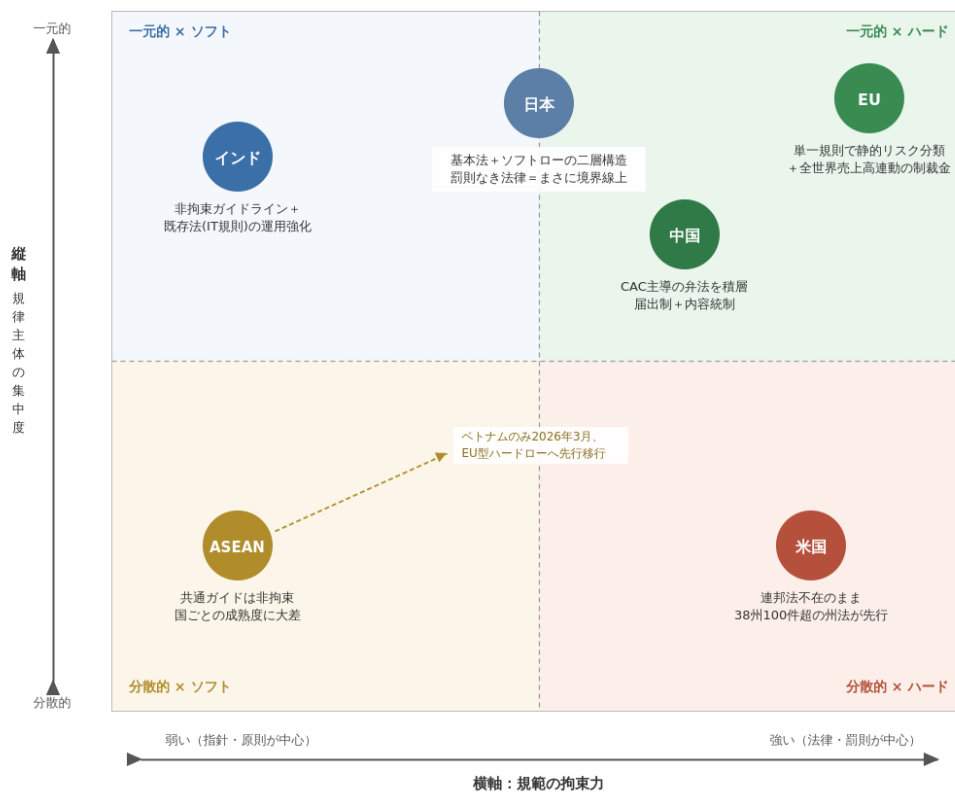
「規範の拘束力は強いのか弱いのか」「規律主体は一元的か分散的か」——この2つの問いを立てるだけで、世界の主要な規制環境の「型」はおおむね見分けられる。以下の表と図1は、この2軸を物差しとして、本号執筆時点(2026年6月)における各法域の現在地を整理したものである。比較の基準点として、第2号で詳述する日本も図中に位置づけた。詳細な制度比較は今後の号で改めて取り上げる。

法域	基本アプローチ	2026年6月時点の特徴
欧州(EU)	ハードロー一本化型	単一の規則(EU AI Act)が、用途を4区分のリスクに静的に分類し、区分ごとに義務・罰則を一括して定める。違反には全世界売上高に連動した高額な制裁金が科され得る。域外適用があるため、EU域内に製品・サービスを提供する日本企業は直接の規律対象となる。詳細は第3号で扱う。
米国	連邦不在・州法パッチワーク型	包括的な連邦AI法は存在せず、カリフォルニア州 SB53(フロンティア AI 透明化法、2026年1月施行)、コロラド州 AI 法(2026年6月施行)、ニューヨーク州 RAISE 法など、38州100件超の州法が先行する。トランプ政権は2025年12月、州法を連邦基準で上書き(プリエンプション)する大統領令に署名し、司法省に AI 訴訟タスクフォースを設置するなど、連邦と州の対立が法廷闘争に発展しつつある。事業者は当面、主要進出州ごとの法規制を個別に監視する必要がある。
中国	国家管理・段階的弁法積層型	全人代レベルの統一AI法はまだ制定されていないが(2026年の立法作業計画で検討加速の方針)、国家インターネット情報弁公室(CAC)主導の行政法規・部門規章が段階的に積み上がる。生成 AI サービス管理暫定弁法(2023年)を土台に、AI 生成合成コンテンツ標識弁法(2025年9月施行、生成コンテンツへの標識義務)、AI 擬人化インタラクティブサービス管理暫定措置(2026年7月施行予定)が相次ぐ。生成 AI サービスは届出・登録制が前提で、2026年4月末時点で868件が届出済み。コンテンツが社会主義の核的価値観に反しないことなど、内容そのものへの統制が欧米とは異なる特徴。中国国内向けにサービスを提供する企業は域外適用の対象となる。
インド	包括法なし・既存法+ソフトロー先行型	EUのような包括的AI法は存在せず、2000年情報技術法など既存法令の適用を基本とする。電子情報技術省(MeitY)が2025年11月に公表した「インド AI ガバナンスガイドライン」は信

		頼・人間中心・規制より革新等7原則を掲げるが、法的拘束力はない。一方、2026年2月施行のIT規則改正は、AI生成コンテンツを「合成生成情報(SGI)」と定義し、プラットフォーム等の仲介者にラベリング・メタデータ埋込みを義務付けるなど、既存の仲介者責任法制を通じた実質的な規律強化が先行している。「規制より革新」を掲げつつ、コンテンツ規律は着実に強化されている点に留意が必要。
東南アジア (ASEAN)	域内非拘束ガイドライン先行・国ごとに温度差	ASEAN が 2024 年に採択した非拘束の「AI ガバナンスと倫理に関するガイド」が域内共通の土台となっているが、各加盟国はそれぞれ異なる速度で制度化を進める「生きた文書」型。シンガポールが最も先行しており、IMDA のモデル AI ガバナンスフレームワークを継続的に改訂し、2026 年 1 月にはエージェント型 AI 向けの新版を公表した。これに対しベトナムは 2026 年 3 月、ASEAN 域内で初めて EU 型のリスク分類を採用した包括的な AI 法を施行予定。インドネシア・タイ・マレーシア・フィリピンなどは、現時点では法的拘束力のないガイドライン段階にとどまる。一国一律のルールを前提とせず、進出国ごとの制度成熟度を個別に確認する必要がある。

規制環境クイックマップー2つの指標で6法域を位置づける

この2軸だけで、世界の主要なAI規制の「型」はおおむね見分けられる（2026年6月時点）



読み方：縦軸の「一元的」は単一の規則・当局に規律が集中する状態。「分散的」は州・国ごとに規律がばらばらな状態。右上に近いほどEU型（統一されたハードロー）、左下に近いほど自主性依存型。
 注：円の位置は本号執筆時点(2026年6月)の定性的な評価であり、各法域の制度は流動的に変化する。詳細は第2号(日本)・第3号(EU)・第4号(米国・中国・グローバルサウス)で扱う。

図1 6法域の規制環境クイックマップ(2026年6月時点)。縦軸＝規律主体の集中度(一元的か分散的か)、横軸＝規範の拘束力(弱いか強い)。右上に近いほどEU型の統一ハードロー、左下に近いほど自主性依存型。比較の基準点として日本も併置した。

図1を2軸で読み直すと、各法域の性格は一層はつきりする。右上(一元的×ハード)にはEUと中国が並ぶが、両者の「ハード」の中身は異なる——EUは議会立法による単一規則と司

法的執行、中国は行政機関(CAC)による弁法の積層と届出制・内容統制であり、企業から見た対応の要点も、EU では事前の適合性確保、中国では当局との継続的な関係管理へと変わる。右下(分散的×ハード)の米国は、拘束力の強い州法が乱立するという意味で、単一のハードローよりもかえって対応コストが高くつき得る象限である。左側(ソフト寄り)のインド・ASEAN も一様ではなく、非拘束の看板の下でコンテンツ規律の実質的強化(インド)やEU 型への先行移行(ベトナム)が進む。そして日本は、罰則を持たない基本法とガイドラインを組み合わせる設計により、拘束力軸のちょうど境界線上に立つ。この見取り図から導かれる実務的な結論は一つである——グローバルに事業を展開する企業は、本社所在地の規制水準を基準にコンプライアンス体制を設計することはできず、進出法域ごとに「型」を見極めたうえで、最も要求水準の高い法域に耐える共通基盤と、法域固有の上乗せ対応とを切り分けて設計する必要がある。この点は、第 4 号(米国・中国・グローバルサウスの動向)で改めて掘り下げる。

読者のための視点：規範を読むときの 3 つの問い

①この規範はハードローかソフトローか(違反した場合に法的な制裁があるか)。②このリスクベースは静的な分類なのか、動的なプロセスなのか。③この規範は技術のどの部分(学習データ、モデル自体、提供方法、利用方法)を対象としているのか。この 3 つの問いを常に立てることで、新しい規制や指針が登場した際にも、本号で示した地図上のどこに位置づければよいかを素早く判断できる。

本シリーズの読み方ガイド

第 2 号以降は、本号で示した地図の上に各規範を位置づけながら読み進めていただきたい。第 2 号は日本の AI 法制(基本法+ソフトローの二層構造)、第 3 号は EU AI Act(ハードロー・リスク分類の典型)、第 4 号は米国・中国・グローバルサウスの動向(ハードロー/ソフトローの分裂と、オープンウェイトモデルが地政学に与える影響)を扱う。第 5 号・第 6 号では視点を転換し、AI を「利用する側」ではなく「開発する側」企業自身の企業統治を扱う。第 7 号以降は、人権・ESG・知的財産・プライバシー・専門職倫理という実体的な論点に進み、第 12 号で全社的なガバナンス体制の構築という形に統合する。

参照すべき一次情報源・規範

- OECD AI 原則(2019 年採択、2024 年改訂)
- G7 広島 AI プロセス国際指針・国際行動規範(2023 年)
- NIST AI Risk Management Framework(Govern/Map/Measure/Manage)
- ISO/IEC 42001:2023(AI マネジメントシステム)
- AI 事業者ガイドライン(総務省・経済産業省、2026 年 6 月時点の最新版は第 1.2 版)
- 人工知能関連技術の研究開発及び活用の推進に関する法律(令和 7 年法律第 53 号)
- EU AI Act(Regulation(EU)2024/1689)

信頼できる文献・解説リソース

6. Vaswani et al., "Attention Is All You Need"(2017 年)——現在の生成 AI の基盤である Transformer アーキテクチャを提示した原論文。技術的な大本を確認したい読者向けの一次文献。
7. Stanford HAI, "AI Index Report"(年次)——AI 技術・投資・規制動向を定量的に整理した報告書。技術と政策の橋渡しとして信頼性が高い。
8. 内閣府 AI 戦略会議公表資料(中間とりまとめ、AI 法案概要等)——日本の AI 法制定過程の一次資料として、立法趣旨を正確に把握するのに有用。
9. OECD.AI Policy Observatory——各国の AI 政策・規制を横断的に検索できる公的データベース。

社内ガイドライン化のポイント

本号で示した技術理解の枠組みと規範の5類型は、社内ガイドラインの「序章・適用範囲」セクションの骨格としてそのまま転用できる。新しい規制・基準やモデルの方式(オープン/クローズド)が登場した際に、まずこの地図上のどこに位置するかを判定する社内手続(規範マッピングプロセス)を定めることを提案する。具体的には、①対象規範がハードローかソフトローか、②静的分類か動的プロセスか、③技術のどの部分を対象とするか、という本号で示した3つの問いを、社内のAIガバナンス委員会(第12号で詳述)が新規規範を評価する際の標準チェック項目として採用することが考えられる。

現時点(2026年)で反映すべき最新動向

2025年9月に全面施行された日本のAI法は、罰則を持たない基本法・理念法であり、AI事業者ガイドライン等のソフトローと組み合わせさせて機能する「基本法+ガイドライン」モデルである点を、地図上の典型例として本号で提示した。同法に基づくAI基本計画は2025年12月23日に閣議決定済みであり、その内容(3原則・4つの基本的な方針)は第2号で詳述する。あわせて、中国発オープンウェイトモデルの普及(詳細は第4号)が、欧米のクローズドモデル中心の競争構造そのものに与えている影響にも触れ、第4号への橋渡しとした。

【注記】本稿の執筆にあたり、AI (Claude, Anthropic) を草稿作成・校正に活用しています。これはAIの強みを執筆のプロセスに組み入れることで、論考の質を高めることを目指した積極的な選択です。当然ながら、内容・法的分析・文責はすべて著者に帰属します。

◆ One Asia Lawyers ◆

「One Asia Lawyers Group」は、アジア全域に展開する日本のクライアントにシームレスで包括的なリーガルアドバイスを提供するために設立された、独立した法律事務所のネットワークです。One Asia Lawyers Groupは、日本・ASEAN・南アジア・オセアニア各国にメンバーファームを有し、各国の法律のスペシャリストで構成され、これら各地域に根差したプラクティカルで、シームレスなリーガルサービスを提供しております。

この記事に関するお問い合わせは、ホームページ <https://oneasia.legal> または info@oneasia.legal までお願いします。

なお、本ニュースレターは、一般的な情報を提供することを目的としたものであり、当グループ・メンバーファームの法的アドバイスを構成するものではなく、また見解に亘る部分は執筆者の個人的見解であり当グループ・メンバーファームの見解ではございません。一般的情報としての性質上、法令の条文や出典の引用を意図的に省略している場合があります。個別具体的事案に係る問題については、必ず各メンバーファーム・弁護士にご相談ください。

◆ アジア ESG/SDGs プラクティスグループ ◆

One Asia Lawyersは、ESG・SDGsと人権DDに関して、東南アジア・南アジア・オセアニアなどの海外においても、各国の法律実務に精通した専門家が、現地に根付いたプラクティカルなアドバイス提供およびニュースレター、セミナーなどを通じて情報発信を行っています。ESG・SDGs・人権DDに関連してご相談がございましたら、以下の各弁護士までお気軽にお問い合わせください。

<著者／アジア ESG/SDGs プラクティスグループ>

	齋藤 彰 One Asia Lawyers Group 顧問 弁護士・神戸大学名誉教授・CEDR 認定調停人 大手海運会社で北米・紅海・欧州向けの自動車専用船の運行管理を経験したのち、研究者への転身を決意。神戸大学法学研究科で比較契約法・国際取引法・国際 ADR 等の教育研究に従事し、学生の国際模擬仲裁大会参加等を促進することにより、法律学のグローバル化に努めてきた。また法科大学院生の海外インターンシップ制度や英語による LL.M.プログラムの創設を主導した。その間に、ICC 仲裁及び調停の実務にも従事し、英国を代表する ADR 機関である CEDR の調停スキルトレーニング (CEDR MST) の日本での初の実施に尽力した。2018 年から One Asia Lawyers の顧問に就任し、実務・教育・研究の架橋に勤めてきた。ビジネスと人権及び海外腐敗慣行防止に向けた規律枠組みの最新動向の調査研究にも取り組んでいる。 akira.saito@oneasia.legal
	難波 泰明 弁護士法人 One Asia 大阪オフィス パートナー弁護士アジア ESG/SDGs プラクティスグループ リーダー 国内の中小企業から上場企業まで幅広い業種の企業の、人事労務、紛争解決、知的財産、倒産処理案件などの企業法務全般を取り扱う。個人の顧客に対しては、労働紛争、交通事故、離婚、相続等の一般民事事件から、インターネット投稿の発信者情報開示、裁判員裁判を含む刑事事件まで幅広く対応。その他、建築瑕疵、追加請負代金請求などの建築紛争、マンション管理に関する理事会、区分所有者からの相談や紛争案件も対応。行政関係では、大阪市債権管理回収アドバイザーを務めるなど、自治体からの債権管理回収に関する個別の相談、研修を担当。包括外部監査人補助者も複数年にわたり務め、活用賞を受賞するなど、自治体実務、監査業務にも精通している。 yasuaki.nanba@oneasia.legal 06-6311-1010
	佐野 和樹 One Asia Lawyers パートナー弁護士（日本法） ミャンマー・マレーシア統括 アジア ESG/SDGs プラクティスグループ 2013 年よりタイで、主に進出支援・登記申請代行・リーガルサポート等を行う M&A Advisory Co., Ltd. で 3 年間勤務。2016 年の One Asia Lawyers 設立時に参画し、ミャンマー事務所・マレーシア事務所にて執務を行う。2019 年にミャンマー人と結婚し、現在はミャンマーに居住しながらミャンマー・マレーシア統括責任者として、アジア法務全般のアドバイスを提供している。 kazuki.sano@oneasia.legal