

Gemara

リスクアセスメント
自動化のための
ガバナンス・リスク・
コンプライアンスの
エンジニアリングモデル



目次

序文.....	3	6. 転換点：センシティブアクティビティ	14
はじめに	3	7. 測定レイヤー	15
1. 適用範囲.....	4	7.1. はじめに.....	15
2. 定義	4	7.2. レイヤー 5：インテントと行動の評価.....	15
3. 基礎概念.....	7	7.3. レイヤー 6：予防措置&是正措置	16
3.1. これまでの取り組み.....	7	7.4. レイヤー 7：監査と継続的監視	17
3.2. リスクアセスメントの現状のアセスメント	7	8. 機械処理に最適化された文書化標準の必要性	18
4. モデル	9	8.1. 機械可読性を超えて.....	18
4.1. モデルの概要.....	9	8.2. 人間のための改善、AI のための改善.....	18
4.2. 完全モデル	9	9. 結論	20
5. 定義レイヤー	11	10. 著者および謝辞.....	20
5.1. はじめに.....	11		
5.2. レイヤー 1：ベクターとガイダンス	11		
5.3. レイヤー 2：脅威とコントロール	12		
5.4. レイヤー 3：リスクとポリシー	13		

序文

ガバナンス、リスク、コンプライアンス（GRC）をソフトウェア開発パイプラインに統合することは、非常に困難な課題です。従来、多くの場合手作業による GRC のアプローチは、現代の開発スピードには適していません。こうした状況から、GRC プロセスに戦略的に工学原理を適用し、より効率的で統合的なプロセスを実現する GRC エンジニアリングという分野が生まれました。その究極の目標は、自動化されたガバナンスを実現することであり、コンプライアンス追跡がデプロイメントパイプライン全体に組み込まれ、コードが本番環境に到達する前に必須の品質ゲートとして機能するようにすることです。

この課題に対する体系的なアプローチの必要性は、CNCF の自動ガバナンス成熟度モデル（AGMM）の作成過程における重要な知見でした。この文書の著者らは、完全自動化されたガバナンスプログラムにおいては、成熟度を少なくとも「ポリシー、アセスメント、執行、監査」という 4 つの異なる領域で測定できると指摘しました。この基礎的な洞察は、セキュアなソフトウェアファクトリーの中核的なアクティビティを記述するための用語集を提供します。Gemara モデルは、業界プロジェクトがこの用語集を適用し始めたことで形作られました。

FINOS Common Cloud Controls (CCC) プロジェクトと OpenSSF の Open Source Project Security (OSPS) Baseline プロジェクトは、AGMM の用語を採用しましたが、より詳細な定義が必要であることを認識しました。そこで、高レベルで抽象的な推奨事項と技術固有の目標を区別するために、さらに 2 つの概念領域を追加しました。このように、4 つの概念から 6 つの領域へと実用的に拡張し、それらが相互に積み重ねられることで、Gemara の階層型モデルが誕生しました。

Gemara モデルは、これらの概念を体系化するために作成されました。自動リスクアセスメントのための GRC エンジニアリングモデルとして、その主な目的は、コンプライアンスアクティビティのカテゴリ、それらの相互作用、およびそれらの間の自動的な相互運用性を実現する方法を記述する論理モデルを提供することです。

はじめに

本稿では、コンプライアンス活動を分類し、それらの機能的な相互作用を定義するために設計された構造である Gemara モデルを紹介します。このモデルは、組織のワークフローや人のプロセスを記述するのではなく、ガバナンスに内在する活動とその構成要素および構造的関係を特定します。これらの活動は実務上長らく存在していましたが、予測可能な交換ポイントを備えた統一的なエンジニアリングアーキテクチャが欠如していました。このモデルは、これらの活動を個別のレイヤーに分解することで、文書化と用語の標準化を促進し、共通リソースの共同保守の基盤を構築します。

1. 適用範囲

本モデルの目的は、ガバナンス、リスク、コンプライアンス（GRC）に関連するアクティビティや、トピックに取り組むための共通基盤を提供することです。同時に、具体的な実装の詳細については、時間の経過とともに発展・成熟していく余地を残しています。

2. 定義

以下に挙げる用語の多くには、業界内で複数の使われ方が存在します。本書では、これらの用語を以下の定義に従って使用します。可能な限り用語について単一の定義を設けることを目指していますが、必ずしもそうならない場合もあり、その際は文脈に基づいてその用法を推測する必要があります。

読者の理解を助けるために、定義された用語は文書全体で頭文字を大文字で表記します。（訳注:本訳では、定義された用語はシングルクォート（' '）で表記します。）また、各セクションにおいて主要な対象として初めて登場する際には太字で表記し、その定義が存在することを示します。以降の記述は、一貫性を保つために大文字で表記します。

- **‘アセスメント’** — (1) 結果が期待された目的を満たしているかを判断するプロセス、または (2) ‘アセスメント’ の中で使用される最小単位のプロセスであり、リソースの ‘コンプライアンス’ が ‘アセスメント要件’ に準拠しているかを判断するためのもの
- **‘アセスメント要件’** — アセスメント者によって満たされていることが確認されるべき、明確で検証可能な条件
- **‘監査’** — 組織のポリシーおよびその状態について、特定の時点で要求事項が満たされていることを検証するために実施される、公式かつ意見を伴うレビュー
- **‘行動評価’** — ミュレーションまたは実世界のアクティビティに対する、意見を伴う観察

- **‘ケイパビリティ’** — システムの機能や特徴であり、攻撃対象領域（アタックサーフェス）を構成する主要要素
- **‘カタログ’** — 関連する記述とメタデータから構成される構造化された集合
- **‘継続的監視’** — 複数のシステムにまたがり、‘アセスメント’ および運用データを継続的に収集するプロセスであり、非コンプライアンスや不正行為の検出、‘是正措置’ の実施、傾向の把握を目的とする
- **‘コントロール’** — (1) システムやリソース、状態に対して望ましい状態を完全に維持する組織の能力、(2) 望ましい状態を実現するためのセーフガードや対策などの手段、(3) ‘目的’ および ‘アセスメント要件’ に関する記述
- **‘コンプライアンス’** — ‘ルール’ または一連の ‘ルール’ に従うこと
- **‘評価’** — ‘アセスメント要件’ に基づき、‘コンプライアンス’ の状態について意見を形成するための手動または自動のプロセス
- **‘エンフォースメント’** — 法令違反の発見とその原因に対応して行われる対応措置
- **‘評価結果’** — アセスメントの証拠およびそれに基づくアセスメント結果

- **‘ガイダンス’** — 関連する‘ベクター’の知識に基づいて、あるトピックや一般的なシナリオにおいて望ましい結果をもたらすことを目的とした文章
- **‘ガイドライン’** — ‘ガイダンス’‘カタログ’を構成する最小単位。多くの場合、最適な実装を設計するための説明的な文脈および推奨事項を含む
- **‘GRC’** — (1) サイバーセキュリティ領域における、‘ガバナンス’、‘リスク’、‘コンプライアンス’分野、(2) 事業部門においてこれらを統合的に遂行するためのプログラム
- **‘ガバナンス’** — 組織およびそのアクティビティに対する戦略的な監督
- **‘インテント評価’** — リソースがポリシーに沿って適切に準備されているか（例：適切な訓練、設定、コード）を確認する‘評価’
- **‘組織’** — 企業、事業部門、チームなど、人的資源、物的資源、仮想資源、情報資源を論理的にグループ化したもの
- **‘脅威’** — 特定の文脈においてベクターの概念が能力に適用され、負の影響を及ぼす‘ケイパビリティ’がある状況または事象
- **‘目的’** — 統一された意図表明であり、複数の状況に応じて適用可能な記述または要件を含む
- **‘オピニオン’** — 評価者の哲学、視点、および能力の制約の中で形成される、現実に対する確信的な近似的判断
- **‘ポリシー’** — 組織の‘リスク許容度’に基づいた、明確に範囲が定義されたルールの集合
- **‘予防措置’** — 非コンプライアンスを引き起こす可能性のあるプロセスを事前に遮断するためのあらゆる措置
- **‘是正措置’** — 実運用中のアクティビティにおいて発生した非コンプライアンスに対する是正措置
- **‘残存リスク’** — ‘リスク軽減’および強制措置の実施後にもなお残る‘リスク’
- **‘リスク’** — 脅威が現実化した際の損失または損害の可能性。組織に対する影響度と発生可能性の組み合わせにより決定される
- **‘リスクカタログ’** — 組織に関連するリスクを体系的にまとめたもの。ポリシーの策定時期や方法を決定するために使用される
- **‘リスク許容度’** — 組織が目標達成のために受け入れる意思のある‘リスク’の水準
- **‘リスクアセスメント’** — システムによって導入される潜在的または実際の‘リスク’を特定するプロセス
- **‘リスク軽減’** — ‘脅威’を防止する、またはその影響を軽減するための対策を策定するプロセス
- **‘リスク受容’** — 低減されていないリスクを必要または不可避なものとして受け入れるという、明確に文書化された意思決定

- **‘ルール’** — 実際に適用・強制可能なポリシー、規制、または法律
- **‘センシティブ アクティビティ’** — 組織にリスクをもたらす種類の行動
- **‘ベクター’** — (1) 攻撃者がシステムの脆弱性を悪用する機会、(2) 不注意により意図しない好ましくない結果を引き起こし得る経路
- **‘脆弱性’** — (1) 本来の用途とは異なる方法で使用された場合に悪用され得る、能力に内在または関連するシステムの弱点、(2) 意図的または非意図的に生じた、**‘コントロール’** の欠如または防御上のギャップであり、害を引き起こすために利用され得るもの

3. 基礎概念

3.1. これまでの取り組み

コンプライアンス マネジメント システム (CMS) の導入を支援するための取り組みは数多く存在します。例えば、ISO19600「Compliance management systems — Guidelines」や、ISO37301「Compliance management systems — Requirements with Guidance for use」などが挙げられます。

また、類似するリソースも数多く存在し、例えば CNCF が最近発表した「Automated Governance Maturity Model」は、組織が自身の状況を測定し、「GRC」プログラムの成長目標を設定するための指標を提供しています。さらに、米国の政府機関である国立標準技術研究所 (NIST) は、[Security Content Automation Protocol \(SCAP\)](#) などの自動化プロトコルの範囲を発展させた、[Open Security Controls Assessment Language \(OSCAL\)](#) に関する取り組みを進めています。

しかしながら、用語、スキーマ、プロセスの一貫した使用には依然として不足があり、また、センシティブアクティビティを統制する際に期待されるべきチーム間の相互運用性も十分とは言えません。このような分断により、組織内および業界全体で手戻り作業が発生し、「ポリシー」と現実の間にギャップが生じています。

セクション 8.1 では、組織がこれらの既存の取り組みからどのように学び、包括的な自動化ソリューションに活かせるかについて説明します。

3.2. リスクアセスメントの現状のアセスメント

‘**リスクアセスメント**’ は一見単純なテーマのように思えるかもしれませんが、実際にはさまざまなシナリオにおいて、多様な形で実施されていることが分かっています。‘**リスク**’ が発生する可能性のあるシナリオを予測するために将来を見据えたアクティビティもあれば、現在実施中または展開中のアクティビティを検証し、‘**リスク**’ が顕在化したかどうか、そしてどのように対応すべきかを判断するために、過去を振り返るアクティビティもあります。

場合によっては、リスクアセスメントは、‘**センシティブ アクティビティ**’ が実行されているライフサイクルと同じタイミングで行われることもありますが、必ずしも明示的に ‘**リスクアセスメント**’ と呼ばれるとは限りません。

これらの ‘**センシティブ アクティビティ**’ は、銀行の窓口係のようなアナログな日常業務から、一見相反するようなソフトウェア開発ライフサイクル (SDLC) まで、さまざまな形態を取り得ます。しかし、結果を ‘**コントロール**’ し、‘**リスク**’ を低減するために実施されるアクティビティの分類は、本質的には共通しています。

セクション 4 で詳しく述べるように、‘**リスクアセスメント**’ は、その入力と出力に基づいて様々な種類に分類できます。これらの分類は累積的であり、それぞれが他の要素を予測可能な形で呼び起こし、層が積み重なっていくように構成されています。

論理的な流れの出発点では、‘**リスク**’ はあくまでも仮説的なものであり、具体的な組織や実施されるアクティビティに関する情報がなければ具体化することはできません。このような初期段階のアクティビティは、コンソーシアムによって実施されたり、‘**ポリシー**’ 作成の過程で日常的に実施されたりします。しかし、その結果としてプロセスが停滞したり、成果物の品質が低下したり、内容が不十分になったりすることがあります。これらの要素が基盤を形成するため、理解が不十分だと、‘**GRC**’ およびセキュリティプログラム全体の有効性が低下する可能性があります。

3.3. 階層的アプローチの確立

‘GRC’ において階層型アーキテクチャモデルを採用する戦略的価値は、懸念事項を明確に分離し、複雑な領域を管理可能な相互接続されたコンポーネントに分解できる点にあります。

Gemaraモデルの構造は、ISO 7498「Information Technology — Open Systems Interconnection — Basic Reference Model」から着想を得ています。一般にOSIモデルとして知られているこのモデルは、ネットワーク機能を明確に分類された階層として記述する、権威あるフレームワークです。

‘GRC’ においても同様のアプローチを採用することで、高レベルの‘ガイダンス’から低レベルの‘エンフォースメント’まで、それぞれのチームやツールが特定のアクティビティに集中しつつ、統合された一貫性の

あるシステムの中で相互運用できるようになります。OSIモデルと同様に、一部のツールは複数のレイヤーにまたがって動作しますが、それでもなお、本モデルに従ってそれぞれのアクティビティを整理・記述することで、明確性の恩恵を受けることができます。

この階層的アプローチを反映して、Gemara モデルは連続した7層からなるアーキテクチャを定義しています。この構造の中で、第4層は重要な転換点となります。すなわち、後続のアセスメントの対象となる形で‘ポリシー’を実装する層です。この層では、ソフトウェア設計やインフラストラクチャのプロビジョニングといったセンシティブ アクティビティが扱われ、要求事項と運用上の現実が交差する橋渡しになります。

4. モデル

4.1. モデルの概要

このアーキテクチャの基本的な原則は、各層がその下の層の上に構築されるという点です。上位層は、下位層の出力やサービスを利用してその機能を実行します。

例えば、レイヤー 5 の‘評価’は、レイヤー 3 で定義された‘ポリシー’への準拠性をテストするように設計されており、その‘ポリシー’はレイヤー 2 の‘コントロール’によって規定されています。このような階層的な依存関係により、‘ガバナンス’ライフサイクル全体を通して、情報と権限の明確な論理的流れが生まれます。

この構造は、従来の文書中心の見方を逆転させます。‘ポリシー’は‘コントロール’の上に構築されます。多くの組織は‘コントロール’を含む‘ポリシー’から始めるため、‘ポリシー’が‘コントロール’の後に続くという考え方は直感に反するよう to 思えるかもしれません。しかし、次の例はモデルを明確に示しています。‘コントロール’が作成されるまで、‘ポリシー’は実装と‘アセスメント’の準備が整いません。

4.2. 完全モデル

このモデルは、定義と測定という 2 つの主要なアクティビティカテゴリに分類されます。これらのカテゴリは、理想的には、それぞれ対応するアクションに関する文書またはログによって表されます。

定義レイヤー（レイヤー 1～レイヤー 3）でのアクティビティはそれぞれ、上位レイヤーまたは自身のレイヤー内で参照可能なドキュメント資産を生成する必要があります。測定レイヤーでのアクティビティ（レイヤー 5～レイヤー 7）の各レイヤーは、それぞれタイムスタンプ付きのログを出力として生成する必要があります。

セクション 3.2 で述べたように、レイヤー 4 は、2 つの半分をつなぐ‘センシティブ アクティビティ’を定義します。最初の 3 つのレイヤーは、‘センシティブ アクティビティ’を指し示し、何が許容され、何が許容されないかを定義します。一方、最後の 3 つのレイヤーは、‘センシティブ アクティビティ’またはその結果を振り返り、定義された期待値に準拠しているかどうかを判断します。レイヤー 7 では、‘監査’によって下位レイヤーの結果の品質を見極めます。

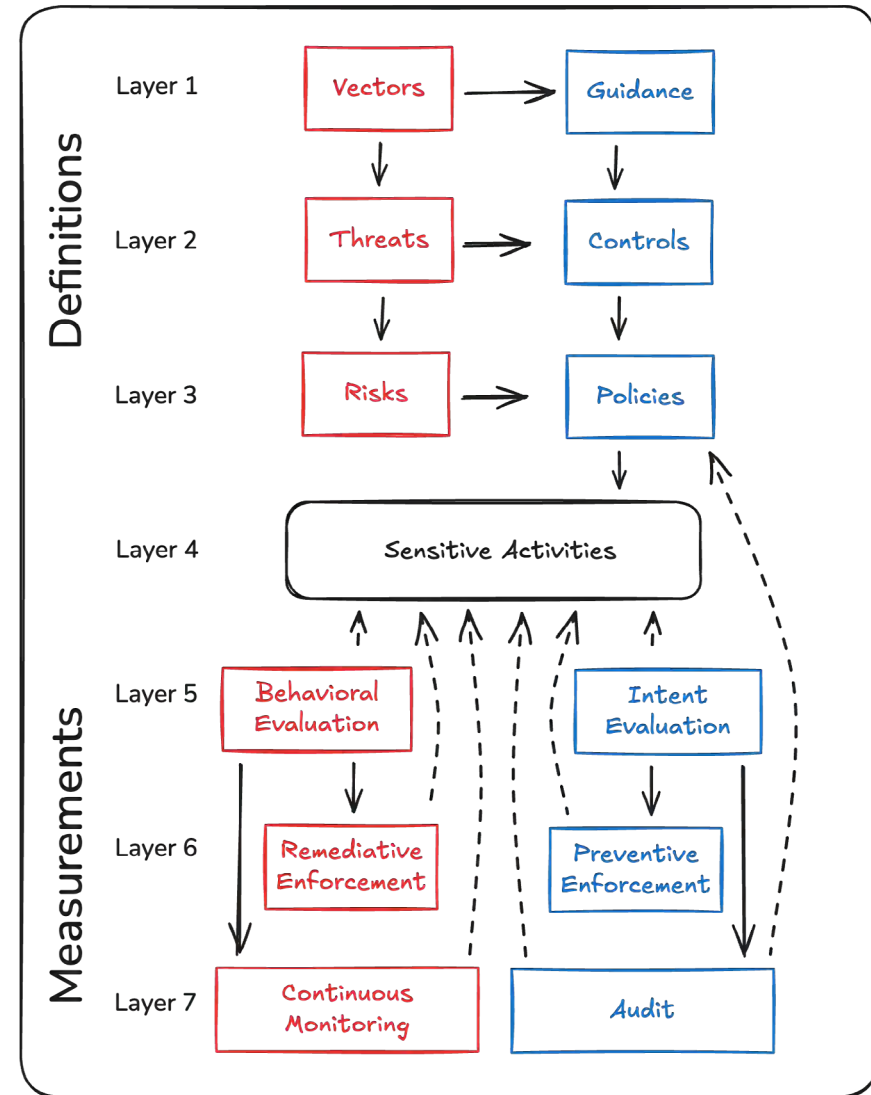
図 1



モデルのアーキテクチャコンポーネントとその要素については、後続のセクションで詳述します。定義レイヤー（セクション 5）から始まり、続いて ‘センシティブ アクティビティ’ を検証する転換点（セクション 6）、最後は測定レイヤー（セクション 7）です。

図2

モデルの関係性と論理フロー



5. 定義レイヤー

5.1. はじめに

第1層から第3層までの‘リスクアセスメント’はすべて連携して機能し、組織が‘センシティブアクティビティ’を実施する際に成功裏に遂行できるよう備えます。

レイヤー1のアクティビティは、好ましくない結果や失敗が発生するさまざまな方法を理解することから始まり、それらの好ましくない結果を防ぐのに役立つ‘ベクター’とそれに対応する‘ガイダンス’の文書化が含まれます。

‘ベクター’を基盤として、レイヤー2では特定のシナリオにおいて‘脅威’がより狭く具体的に定義されます。これらの‘脅威’は、‘コントロール’の正当性を文書化するものであり、‘コントロール’は、関係者が‘脅威’を軽減するための明確な目的と要件を提供します。

強固なセキュリティ実装を実現するための総仕上げとして、レイヤー3は組織が直面する‘リスク’に優先順位を付け、怠慢、ミス、悪意のある行為につながる最も切迫した機会を軽減するために必要な‘ポリシー’を定義します。

これら3つの層が効率的な方法で連携されると、実装アクティビティを加速させ、強化するための設計要件として機能します。

しかし、現実にはその逆のケースの方が多く見られます。準備や設計要件の一部として定義を適切に調整しないと、‘コンプライアンス’がセキュリティから切り離されてしまうという事態が必然的に発生します。‘コンプライアンス’アクティビティをより大きなイニシアチブの戦略的な一部として活用するのではなく、‘コンプライアンス’自体が目標となってしまうのです。

「指標が目標になると、それはもはや良い指標ではなくなる」と [GoodHart's law](#) が教えてくれるように、‘コンプライアンス’のための‘コンプライアンス’を義務付けると、私たち自身の有効性が低下してしまいます。そのため、理想的なセキュリティを実現するには、これら3つのレイヤーの背景にある論理を深く理解することが不可欠です。

5.2. レイヤー1: ‘ベクター’ と ‘ガイダンス’

包括的で高水準な‘リスクアセスメント’の必要性は、通常、アセスメント対象となるアクティビティの範囲からは大きく離れた要因や要件によって明らかになります。法律などの‘ルール’によって必要性が明確になる場合もあれば、人工知能の出現のように、まだ十分にアセスメントされていない新しい技術分野における‘コントロール’の前提条件として、この種のアクティビティが求められる場合もあります。

‘ベクター’を文書化する際、各ステップで使用される技術など、技術的な詳細を理解する必要はありません。むしろ、誤りや悪意が生じる可能性に焦点を当てます。これらは個別に文書化することも、‘カタログ’内にまとめることもできます。また、独立した成果物として公開することも、関連する‘ガイダンス’と併せて公開することも可能です。‘ベクター’の例は、[MITRE ATT&CK](#) フレームワークのテクニックとして見つけることができます。

‘ガイダンス’を構成する各要素（‘ガイドライン’と呼ばれる）は、通常、単独で存在するものではなく、多くの場合、長期にわたる‘ガイダンスカタログ’として発行されます。各‘ガイドライン’には、実施の詳細を事前に知らなくても最適な結果を設計するための説明的な背景情報や推奨事項が含まれていることがよくあります。

‘ガイダンス’は、特殊な状況に合わせて社内で作成される場合もありますが、多くの場合、業界団体、政府機関、または国際標準化団体によって策定されます。例としては、[OWASP Top 10](#), [NIST Cybersecurity](#)

Framework, HIPAA, GDPR, CRA, または PCI および ISO 規格などが挙げられます。

図 3 に示すように、‘ベクター’ アーティファクトは作成作業を迅速化し精度を高めるために、‘ガイダンス’ と ‘脅威’ の両方から参照できます。同様に、‘ガイダンス’ アーティファクトは、特定の ‘コントロール’ がそれぞれの ‘ガイドライン’ をどのように適用するかを示すために、‘コントロール’ から参照できます。

5.3. レイヤー 2: ‘脅威’ と ‘コントロール’

‘ベクター’ と ‘ガイダンス’ を作成する高レベルのアセスメントに基づいて、組織やコンソーシアムは、依然として仮説的な ‘リスクアセスメント’ を頻繁に実施します。すなわち、‘リスクカタログ’ を考慮せず、代わりに明確な ‘アセスメント要件’ によって裏付けられた ‘目的’ を持つ、技術固有の脅威情報に基づく ‘コントロール’ に焦点を当てます。レイヤー 1 で説明したアクティビティと同様に、これらは仮説的な ‘リスクアセスメント’ であり、後で組織の ‘リスクカタログ’ のコンテキストで完全に運用化することができます。

情報技術のエコシステムにおいて、‘脅威’ と ‘コントロール’ という用語が頻繁に過剰に使用されることで曖昧さが生じています。汎用的な資産と特定の資産が作成される場合、それらは両方とも ‘コントロール’ と呼ばれることが多く、結果として ‘ポリシー’ にばらつきが生じ、品質に一貫性がなくなります。セクション 2 で説明する ‘脅威’ と ‘コントロール’ の定義は、これらの用語が使用される様々な方法の中から選ばれたものです。

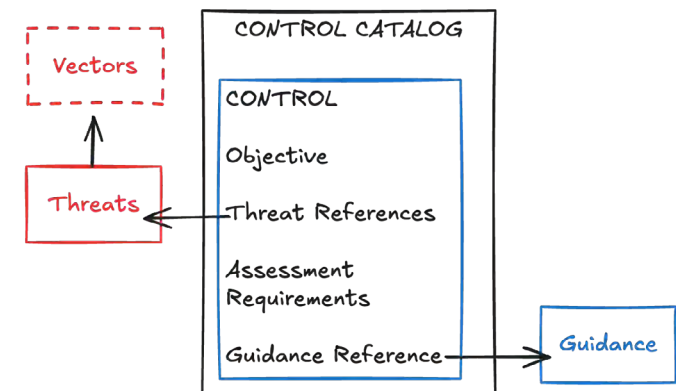
理想的には、‘脅威’ は既知の攻撃経路から抽出され、‘ケイパビリティ’、場所、行動などの特定の要素にマッピングされることで、‘脅威’ がいつ、どこで発生する可能性があるかが明確になります。‘脅威’ は特定のシナリオ、事業部門、またはテクノロジーに関連付けられる場合があり、組織の ‘ポリシー’ 策定に役立ちます。このようにして、対応する要素が存在しなくなった場合には、‘コントロール’ を除外対象としてマークすることができます。

‘コントロール’ は、確立された ‘ガイダンス’ に基づいて策定され、アセスメント者が直接実施できる ‘アセスメント要件’ を含むべきです。これらは通常、組織が内部利用のために策定、あるいは業界団体、政府機関、または国際標準化団体が汎用的に策定します。

例としては、CIS Benchmarks, FINOS CCC, OSPS Baseline などが挙げられます。

従来、‘コントロール’ は ‘コントロールカタログ’ としてトピックごとにまとめて提供されていましたが、情報システム向けに再利用可能な汎用的な ‘コントロール’ を定義するコンポーザビリティへの傾向が強まっています。これらは、一般的な記述という点では ‘ガイダンス’ に似ていますが、レイヤー 2 の資産として機能するのに十分な詳細な ‘アセスメント要件’ を備えています。その一例として、FINOS Common Cloud Controls (FINOS CCC) の「コアカタログ」が挙げられます。このカタログには、‘脅威’ と ‘コントロール’ の両方の定義が含まれており、これらは後でテクノロジー固有のカタログにインポートされます。

図3



情報システムの‘コントロールカタログ’を作成するための推奨手順は、まず技術の機能をアセスメントし、次にそれらの機能に対する‘脅威’を特定し、最後にそれらの‘脅威’を軽減するための‘コントロール’を開発します。しかし、組織がまず作成者の経験に基づいて‘コントロール’を作成し、その後で‘脅威’を追加して、それらの‘コントロール’の正当性や根拠を強化することは珍しくありません。

5.4. レイヤー 3: ‘リスク’ と ‘ポリシー’

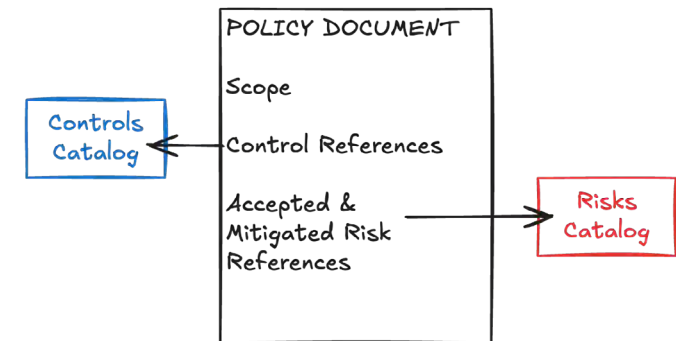
あらゆる規制や指針のそれぞれの‘コントロール’は、組織がそれをいつ、どのように参照するかによって厳密に左右されます。これは、‘リスクアセスメント’が概念段階を終え、現実と向き合う段階です。適切に実施されれば、最も関連性の高い情報に基づいて、適切な‘コントロール’が選択されることが保証されます。

‘リスク’を正確に定義するには、組織の置かれている状況、つまり技術的な詳細から地政学的な詳細まで、あらゆる側面をしっかりと理解する必要があります。通常、‘リスクアセスメント’は‘脅威アセスメント’に基づいて行われ、好ましくない結果が生じる状況的可能性を算出し、それを特定の‘センシティブアクティビティ’の状況下で組織に及ぼす影響と比較します。これには、使用される技術、関与するスタッフ、および取り扱われる資源の価値などの要素が含まれる可能性があります。

しかし、すべての‘リスク’が同じように扱われるわけではありません。‘**リスク許容度**’とは、組織が目標達成のために受け入れる‘リスク’のレベルです。‘リスク許容度’は、‘リスクカタログ’に明確に示され、組織のルールをいつ、どのように作成するかを決定するために使用されます。‘ポリシー’は、組織の‘リスク許容度’に基づいて明確に範囲が定められた一連のルールです。‘ポリシー’は、ベストプラクティスと業界標準に基づいているものの、組織に合わせて調整されたガバナンスルールを提供します。‘ポリシー’は必然的にある程度の‘**リスク受容**’

を伴うため、組織固有の‘リスク許容度’を考慮せずに、それらを適切に開発することはできません。

図4



完全な‘ポリシー’文書には期限が定められ、‘アセスメント要件’付き‘脅威’情報に基づく‘コントロール’への参照が含まれ、さらに、その影響を受ける関係者に展開するための明確な計画が含まれます。

‘ポリシー’文書は他の‘ポリシー’文書から参照されることがあり、機能的な継承モデルを形成します。これらの文書は、設計要件として関係者に配布され、レイヤー 5 アセスメントの出発点として使用できます。‘リスク’を‘ポリシー’の配布に含め、‘リスク’を特定の‘脅威’とそれに続く‘コントロール’にマッピングすることで、実際に成功する可能性が大幅に高まります。

計画段階で作成された‘ポリシー’文書は、セキュリティが後々障害となるのではなく、‘センシティブアクティビティ’に組み込まれるようにするための機能設計要件として機能します。

6. 転換点：センシティブアクティビティ

‘センシティブアクティビティ’とは、組織に‘リスク’をもたらす可能性のあるあらゆる行動を指し、ガバナンスの必要性を生み出します。これらは Gemara モデルの中核を成すものであり、歴史を通じて GRC が様々な形で存在してきた根本的な理由でもあります。

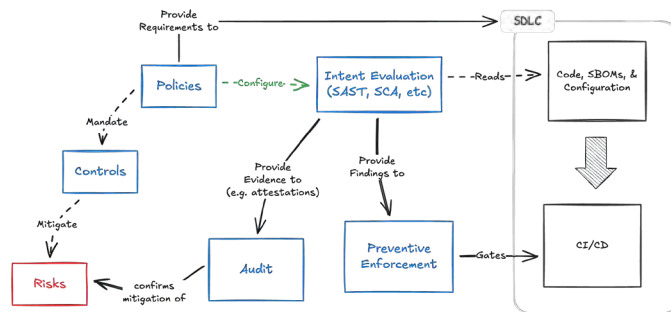
‘センシティブアクティビティ’は、さまざまな形態をとる可能性があります。例えば、銀行支店の建設方法を規制する建築基準、銀行員が法律を遵守するために従わなければならない規則、開発チームが新しいモバイルアプリを作成する際に従うべきプロセス、そしてソフトウェアコード自体に課せられる要件などです。

ツールによってはこのプロセスの複数の要素を提供していますが、リーダーはモデルを参照することで、すべての重要なアクティビティが実施されていることを確認できます。

図 6 では、より詳細な例を示し、ライブ実行環境のアクセスメントへとプロセスを継続しています。完全に統合されると、評価結果はソフトウェア開発ライフサイクルの最終段階として機能します。これは、品質保証、ユーザー受け入れテスト、製品フィードバックと同様です。

図5

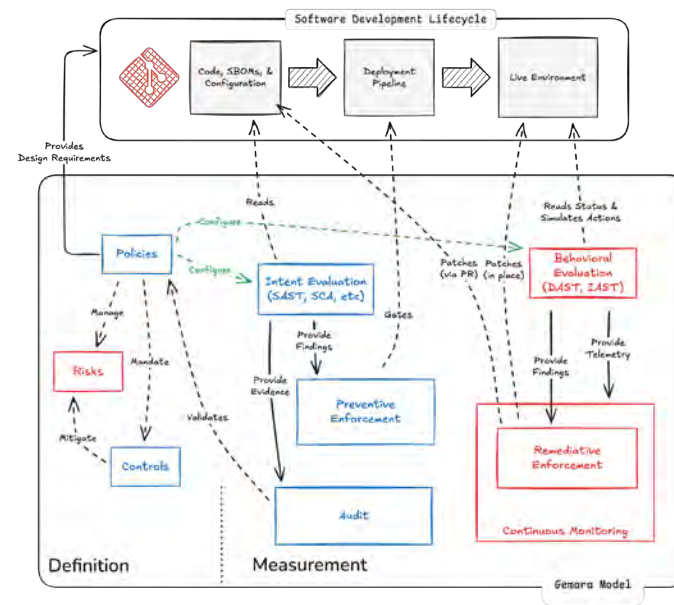
サプライチェーンへのモデルの適用安全



このモデルは、既に一般的に行われている作業に不可欠なコンテキストを提供するのに役立ちます。図 5 では、このモデルを用いて、安定した安全なソフトウェアサプライチェーンを確保するために必要なさまざまな作業カテゴリを示しています。ポリシーは、開発ワークフローと、ソフトウェア構成分析（SCA）などのアセスメントツールの両方に情報を提供するために使用されます。完全な成熟度を実現するには、SCA アセスメント結果を、コンプライアンス監査だけでなく、執行メカニズムにも統合する必要があります。

図6

モデルを完全なものに適用するソフトウェア開発ライフサイクル



7. 測定レイヤー

7.1. はじめに

下位層で行われる準備作業によってセンシティブアクティビティが安全に計画・実行されることを確実にするのに加え、最終の3つの層では結果を振り返り、組織の‘ポリシー’が遵守されていることを確認します。

意図された結果と実際の結果を検証することから始まる第5層のアクティビティは、‘インテント評価’と‘行動評価’のいずれかに分類できます。

‘評価結果’に基づき、第6層では、ガードレールとして機能する‘予防措置’（準拠していない設計が稼働する前に阻止する）と、現実世界のシナリオで悪影響が検出された後に修正を行う‘是正措置’アクティビティについて説明します。

最後に、レイヤー7では、組織のポリシー、評価、および強制措置の有効性を監査し、センシティブアクティビティが永続的にコンプライアンスを遵守し続けることを保証するための継続的なモニタリングを調整するアクティビティについて説明します。

古くから伝わるリーダーシップの格言に、「部隊は指揮官がしっかりと検査したものだけをうまく遂行する」というものがあります。第5層、第6層、第7層に分類されるアクティビティは、まさにその検査の役割を果たし、組織が卓越した成果を上げるための基盤となります。

7.2. レイヤー5：インテントと行動の評価

レイヤー3の‘リスクアセスメント’アクティビティがポリシーを生成するのと同様に、レイヤー5のアクティビティはポリシーコンプライアンスの評価を通じて意見を生成します。これらの‘評価’は、‘インテント評価’または‘行動評価’いずれかの形式を取る構造化された検査で構成されます。

これら2種類の‘評価’は必ずしも並行して行う必要はありませんが、結果をまとめてから実施判断を行う方が、強制措置を行う上で最も有益な情報が得られます。

評価は通常、独立した故障モードを持つ複数の‘アセスメント’で構成されます。組織の方針へのコンプライアンスを確保し、軽減されていない‘リスク’をアセスメントする最善の方法は、各‘アセスメント’と、特定の管理策に対して定義された‘アセスメント要件’との関係を明確に文書化することです。

‘インテント評価’の例には、人間同士のインタビューと、構成の自動化または手動による検査の両方が含まれます。情報システムでは、従来の監査アクティビティの多くはダッシュボードのスクリーンショットなどの手動による‘インテント評価’に依存していましたが、最近の傾向は自動化された構成スキャンとコード分析へとますます移行しています。

アナログ的な‘インテント評価’では、チームの手順、トレーニング記録、ハードウェア、設備などが対象となる場合があります。デジタル版では、ソフトウェア構成分析やクラウドリソース構成のスキャンなどがこれに該当します。

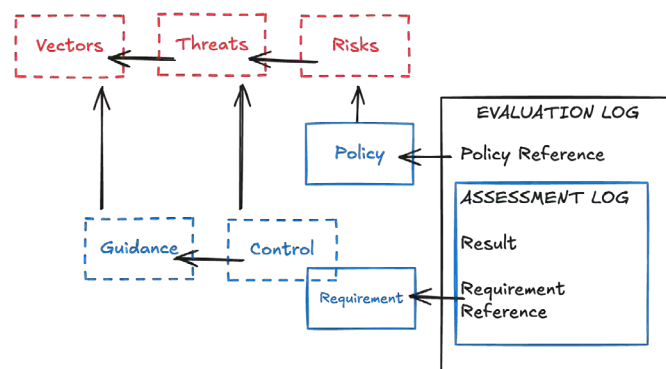
‘行動評価’は、期待される結果が達成されることを確認するためにユーザーの行動を観察またはシミュレートするあらゆる行動の形をとる可能性があります。悪い行動のシミュレートはセキュリティ上の欠陥や‘ポリシー’違反を特定するのに役立つだけでなく、適切な動作を定期的にシミュレーションすることは、システムが常に期待どおりに動作することを保証する上でも有効です。アナログ環境では、これは覆面調査員による調査として現れる可能性があり、情報システムの評価にはペネトレーションテストが含まれる場合があります。

評価者の‘オピニオン’は、その評価者の哲学、視点、能力（あるいは評価者自身の能力）の制約の中で形成された、現実に対する確固たる近似値です。これは事後的な‘リスクアセスメント’であり、どのような‘リスク’が導入または提案されたかを理解するためのものです。

このプロセスでは、評価結果を組織の期待値と比較します。組織の期待値には、理想的には、‘ポリシー’、‘コントロール’、および‘アセスメント要件’に反映されます。評価者はベンダーや業界団体から提供される場合もありますが、堅牢な‘評価’には、‘コンプライアンス’プログラムのニーズに合わせて‘アセスメント’をカスタマイズするために、特定の‘ポリシー’に基づいて行われるべきです。

‘図7’に示すように、各成果物間の関係（「マッピング」とも呼ばれる）を適切に維持することで、‘評価’は、システムが様々な関連成果物の‘コンプライアンス’状態を示すための複数の機会を提供することができます。

図7
評価ログのマッピング先
定義アーティファクト



7.3. レイヤー 6：予防措置&是正措置

‘評価’で不適合事項が記録された場合、次のステップは‘エンフォースメント’を講じることです。これは多くの場合、‘評価結果’が特定されたのと同じツールまたはプロセス内で行われ、アラート、中断、または是正が実行されます。後者の2つは‘エンフォースメント’の行動とみなされ、アラートなどの継続的なアクティビティは‘継続的監視’の一部です。

このアクティビティはしばしば純粋に客観的なものと見なされがちですが、私たちはこれを‘リスクアセスメント’の継続として捉えています。なぜなら、執行段階では、法令違反が判明した場合に取るべき最善の行動方針についてオピニオンを形成する必要があるからです。予防策を講じるか、あるいは修復策を講じるかを選択する際には、それぞれの行動方針に伴う‘リスク’を理解することが不可欠です。

‘予防措置’とは、法令違反を防止するためにプロセスを中断するあらゆるエンフォースメントを指します。これは主に、不適切な設定などの悪意のある行為への対応として行われますが、センシティブアクティビティの展開前に実行される、より複雑な‘行動評価’プロセスも含まれる場合があります。

ソフトウェア開発ライフサイクルにおいては、これはデプロイメントゲートという形をとることがあり、物理的な施設における境界ゲートのように、入退室を制御する役割を果たします。これにより、意図的に行われるすべての行動がポリシーに従って審査されることが保証されます。

‘是正措置’とは、展開されたアクティビティにおけるコンプライアンス違反への対応するためのエンフォースメントを指します。これには、コンプライアンス違反のリソースを隔離、置換、再トレーニング、またはコンプライアンスに準拠した構成で再展開することが含まれる場合があります。通常、これには、対象の状態を検出するための何らかの監視機能と、環境への変更、または定義された運用しきい値に基づいてポリシーに沿った修復を確実に行うための最新の適用ポリシーとの連携が必要です。

多くの修復作業は不安定な性質を持ち、エンドユーザーが混乱する可能性もあるため、修復ツールと併せて適切なログ記録とアラート機能を必ず導入する必要があります。

図7で述べた点をさらに掘り下げると、‘エンフォースメント’アクティビティは‘評価’の結果に適切に結び付けられることで、より高い効果を発揮する可能性があります。これにより、‘エンフォースメント’が一方的に適用されるのではなく、変更に関わるすべての人やシステムがその正当性を完全に理解できるようになります。

7.4. レイヤー7: ‘監査’ と ‘継続的監視’

‘コンプライアンス’においておそらく最も有名で、同時に最も恐れられている部分は、センシティブアクティビティに関与していない人々にコンプライアンスを証明するプロセスです。

‘監査’とは、組織のポリシーと姿勢を正式かつ客観的にレビューし、過去の‘リスクアセスメント’の質と‘残存リスク’を評価するものです。

‘監査’は、ある時点における‘コンプライアンス’状況の近似値であるため、その結果は、入手可能な証拠や収集された事実の観察に基づく確固たる主張となります。‘監査’の範囲は多岐にわたり、組織が遵守する‘ガイダンス’、策定した‘コントロール’、実施している‘ポリシー’、‘評価’方法、および‘エンフォースメント’状況に関する意見形成が含

まれる場合があります。意見は通常、地政学的な規制要件と、当該業界および関連リソースにおける既知のベストプラクティスを組み合わせた情報に基づいて形成されます。

‘監査’は通常、まとめて実施され、‘評価’アクティビティ中に収集された特定の時点における証拠を探します。従来、‘監査’ではその簡便性と予測可能性から手動による‘インテント評価’が好まれてきたが、サイバーセキュリティ監査では、‘インテント評価’と‘行動評価’の両方において自動化ツールの導入がますます進んでいます。

監査アクティビティは過去の検査に重点を置くことが多いのに対し、‘継続的監視’は評価結果と運用指標をリアルタイムで集約します。このプロセスは、‘脅威’や‘コントロール’の不備が発生した時点で特定し、‘是正措置’の実施を支援します。これは「CCM」(継続的‘コンプライアンス’監視)とも呼ばれます。

成熟したCCM運用では、‘ポリシー’に基づいた継続的なプロセスを確立し、複数のシステムを活用することで、展開されているすべての資産の可視性を常に最大化します。これは物理的な例で言えば、継続的なセキュリティ巡回、ライブ映像監視、そして稼働中の周辺警備といったものに相当します。

8. 機械処理に最適化された文書化標準の必要性

8.1. 機械可読性を超えて

セクション 3.1 で述べたように、リスクアセスメント自動化のためのガバナンス・リスク・コンプライアンスのエンジニアリングモデルは、これまでの多くの反復と学びの成果です。

OSCAL に関する取り組みからは計り知れない価値が生まれており、そこでは、厳しく規制された組織によって、公的にも私的にも、機械可読な文書が大量に作成されてきました。OSCAL 分野で行われた取り組みは、数え切れないほどの企業に大きな利益をもたらしてきたため、決して見過ごしたり軽視したりすることはできません。

しかし、機械可読な「GRC」文書を作成する際に、機械向けに最適化することの利点については言及すべき点があります。関連するスキーマを作成するこれまでの取り組みでは、2 つの重要な点で期待に応えられませんでした。

まず、最も明白な欠点は、「ポリシー」という用語をセキュリティツールの構成を説明するために誤用しているツールにあります。これらのリソースは確かに有用ですが、多くの組織は、Policy-as-Code を作成するという名目で、高度な訓練を受けた GRC およびセキュリティ専門家にツール構成の管理を任せることで、彼らの潜在能力を無駄にしています。

2 つ目の欠点は発見しにくいもので、カスタマイズ可能なフィールドを許容する仕様という形で現れます。

無害にあるいは必要にさえ聞こえるかもしれませんが、GRC 文書スキーマにカスタマイズを追加すると、これらの分野に関する事前知識を持つカスタムツールが必要になります。この柔軟性の代償として、データまわりで広く採用されるツールを構築するのに必要な構造的明確さが失われてしまいます。スキーマはさまざまな用途をサポートするのに十分な柔軟性を持つべきですが、また、自動化されたエコシステムがその周りに立ち上がり繁栄することを可能にするだけの明確な意見も持つべきものです。

これまでの成功事例から得た教訓を活かし、既知の課題に対処することで、組織は複雑な問題に対処できるほど堅牢でありながら、あらゆる専門家が迅速に理解できるほど分かりやすく、さらに OSCAL などの他のフォーマットへのマッピング機能も維持できる、非常に効率的な GRC エコシステムを構築できるようになります。

各アクティビティタイプごとに、明確な意見に基づいた標準化されたスキーマを確立することで、業界全体で自動化された「リスクアセスメント」を迅速に加速させることが可能になります。

8.2. 人間のための改善、AI のための改善

最適化された「GRC」エンジニアリング戦略を採用する組織は、人間と人工知能の両方にとってメリットを得ることができます。

コミュニケーションと引き継ぎは、規制の厳しい企業において障害となることがよく知られています。「ポリシー」を関係者に配布すること、変更を管理すること、意図しない手戻り、さらには誤解による不必要な作業に数ヶ月を費やすことなどの課題が常に存在することが、フィールド観察で明らかになっています。

文書スキーマを標準化する企業は、文書を適切な場所に適切なタイミングで届けるために必要なツールを開発できます。こうした通信最適化を構築するために必要なデータベースと API は、二次的なメリットももたらします。それは、モデルコンテキストプロトコル (MCP) の基盤となることです。

MCP や AI スキルといった適切に設計されたシステムを、機械学習に最適化された GRC データベースに重ね合わせると、曖昧で構造化されていない文章を、固有の技術識別子と明確なリソースマッピングに置き換えることで、AI の機能が強化されます。構造が明確になることによって、モデルは要件に基づいて動作できるようになり、意図の確率的推論の必要性が最小限に抑えられ、より決定論的な結果が得られます。

GRC（ガバナンス、‘リスク’ 管理、コンプライアンス）に特化したシステムが、エンジニアリングに特化したシステムに、ほぼ決定論的な方法で正確なコンテキスト情報を提供できる Agent2Agent (A2A) 対応システムと連携することで、AI システムが必要とするリソースが大幅に削減され、セキュリティツールの精度が向上し、誤検知や誤検出に費やす時間が減り、組織に真の価値を提供するための時間が増えることが観察されています。

9. 結論

このモデルは、最新の‘ガバナンス’、‘リスク’、‘コンプライアンス’プログラムを構築するための、明確で実用的かつ拡張可能なフレームワークを提供します。共通の用語と論理的で階層的なアーキテクチャを確立することで、高レベルの‘ガイダンス’、技術的な‘コントロール’、組織の‘ポリシー’、自動化された‘エンフォースメント’といった要素間の複雑な相互作用を分かりやすく解説します。

Gemara プロジェクトは Linux Foundation 内で設立され、Open Source Security Foundation (OpenSSF) が管理しています。このプロジェクトは、本モデルに記載されているすべてのアクティビティタイプに対応した、機械処理に最適化されたドキュメントスキーマ、役立つ用語集、およびソフトウェア開発キット (SDK) の維持管理を担当しています。プロジェクトのガバナンス構造により、メンテナンスは複数の組織に完全に分散されており、Sonatype、Red Hat、CVS Health が創設時のメンテナーを務めています。

コミュニティが維持管理するスキーマは、拡張可能であり、場合によってはコミュニティ全体を最適にサポートするために修正も可能です。同様にオープンソースの SDK は、Gemara 互換ドキュメントの読み取りと変換を支援するように設計されています。これらのリソースにより、あらゆる組織が迅速に社内の‘GRC’エンジニアリングソリューションを立ち上げることができ、エコシステムを加速するソリューション構築を目指す人は誰でも数分で作業を開始できます。

Gemara プロジェクトは、業界全体の高度なフレームワークと技術固有のセキュリティ対策との間のギャップを埋めることで、自動化されたガバナンスへの体系的な道筋を提供します。‘GRC’エンジニアリングのための共通の論理的基盤を提供することで、このプロジェクトは抽象的な概念を超え、現代のセキュアなソフトウェアファクトリーのための、実用的で機械最適化された基盤を提供します。

10. 著者および謝辞

本書全体を通して述べられているように、本書の内容は無数の教訓や情報源から引用されており、その全体像を包括的に特定することは困難です。以下に、これらの知見をこのモデルにまとめた著者、および最終的な形に影響を与え、あるいは影響するガイドを与えた主要な貢献者の一覧を示します。

著者：Eddie Knight, Jennifer Power

寄稿者：Hannah Braswell, Travis Truman, Brandt Keller, Damien Burks, Sonali Mendis, Michael Lysaght, Robert Griffiths, Rob Moffat, Naseer Mohammad, Stevie Shellis, Jared Lambert, Andrew Martin, Sonu Preetam, Stephanie Harris, Marcus Burghardt, Sally Cooper

顕著な影響を与えた方：Ian Miell, Michaela Iorga, Tara Houlden, Ben Cotton, Christopher Robinson, Evan Anderson, Maxime Coquerel, David Wheeler, Kamran Kazmi, Ian Tivey, Dana Wang, Jeff Diecks

ありがとうございます！
ぜひご参加ください。
gemara.openssf.org

この資料は、OpenSSF 発行の「[Gemara: A Governance, Risk, and Compliance Engineering Model for Automated Risk Assessment](#)」を OpenSSF Japan Chapter の有志メンバーで日本語に翻訳したものです。

日本語版翻訳協力：OpenSSF Japan Chapter 翻訳チーム

- ・清海 佑太（本田技術研究所）
- ・川名 のん（日立製作所）
- ・下沢 拓（日立製作所）
- ・余保 束（ルネサスエレクトロニクス）
- ・池田 宗広（サイバートラスト）

