

最初から
セキュアなAIを
構築することが
重要な理由





人工知能（AI）プロジェクトの設計者が、設計の初期段階から強固なセキュリティとデータガバナンスを組み込まずに取り組みを始めてしまうと、サイバーリスクを増大させてしまう恐れがあります。ランド研究所の国家安全保障調査部門の調査によると、生成AIプロジェクトは、従来のITプロジェクトに比べて失敗率が高く、最大で5件中4件が失敗に至る可能性があるとの見解が示されています。その原因の多くは、脆弱性への未対策、性急な導入、そしてAI機能とアプリケーションの性急な開発リリースサイクルにあるとされています。

では、企業はどのようにAIの安全性を初期段階から確保し、プロジェクトを成功に導けばよいのでしょうか？

その答えはプロアクティブでセキュリティを第一に考えたアプローチにあると、ベライゾンでは考えています。高度なサイバーセキュリティとネットワークソリューションを初期段階から統合することで、リスクを軽減し、生成AIの可能性を責任ある形で引き出すことができるのです。

成長する生成AI

生成AIはアジア太平洋地域（APAC）全体で産業の変革を加速させており、IDCの試算によれば、この地域のAI市場規模は2027年までに260億ドルに達する可能性があると考えられています。また、ある業界のインサイトによるとオーストラリアでは企業の約70%がすでにAIを活用しており、85%以上が今年中に生成AIへの投資をさらに拡大する予定だということです。この急速な導入はイノベーションを促進する一方、適切なセキュリティ対策がなければ、組織はサイバー脅威にさらされる可能性があります。

米国国立標準技術研究所（NIST）は、AIアルゴリズムのセキュリティ確保は依然として複雑な課題であり、いくつかの問題が未解決であると指摘しています。さらに、Economist Impactのレポートでは、オーストラリアの企業の47%が現在のデータおよびAIガバナンスの枠組みにギャップがあると認識していると強調しており、ベライゾンの専門知識こそがそのギャップを埋めるのに役立つのです。

安全な一歩を踏み出すために

ベライゾンの2025年度データ漏洩/侵害調査報告書（DBIR）は、ランサムウェアの被害組織の54%のドメインが認証情報ダンプ内に表示されていることを浮き彫りにし、AIシステムでは最初から認証情報を保護することが非常に重要であると強調しています。自動化された脆弱性管理などのベライゾンのソリューションは、これらのリスクを迅速に軽減するのに役立ちます。

ベライゾンのアプローチは、既知の脆弱性を迅速に修正することから始まります。

最近ソウルで開催されたイベントで、ベライゾンの情報セキュリティのグローバル責任者であるAlistair Nealは、特に「サービスとしてのランサムウェア」を提供する「RansomHub」プラットフォームのような新しい攻撃戦略の進化に伴い、増加する脅威の中でAIセキュリティの自動化がどのように脆弱性管理とガバナンスを強化するかについて解説しました。

Alistair Nealは、AIセキュリティの自動化とセキュリティ環境の変化が、企業にデータ共有やコラボレーションを拡大する絶好の機会を提供していると指摘しました。一方で、データ保護とセキュリティ管理には体系的なアプローチが不可欠なのです。ベライゾンのプロアクティブな監視および対応能力は、企業がそのような危険に先手を打つことをサポートし、プロジェクトが拡大する前から生成AIインフラを守ることを可能にします。

ネットワークの変革も同様に重要です。生成AIの需要が高まる中、組織には堅牢で拡張性の高いネットワークが求められています。次世代プラットフォームを導入したとしても、攻撃対象領域の拡大はセキュリティの責任者にとって課題となります。ベライゾンのプライベート5Gネットワークは、こうした変革をサポートし、AIワークロードに対応するための大容量の接続を提供します。また、ベライゾンのセキュアゲートウェイはデータフローを暗号化し、生成AIシステムがテスト段階から本格導入に移行する際の安全性を確保します。

NISTの「敵対的機械学習」に関するレポート（NIST.AI.100-2）の実例では、自動運転車の生成AIモデルにおいて、標識分類を守るためのトレーニングの安全対策が不十分だったため、破損した一時停止の標識を速度制限と誤読し、攻撃者に悪用される脆弱性が生じたケースが紹介されました。ベライゾンのサイバーセキュリティソリューションは、開発者がこのような機密性の高いトレーニングデータをテストし、保護するのに役立ちます。

ベライゾンとの包括的なアプローチの構築

生成AIのセキュリティを確保するには、リスク管理、データガバナンス、サイバーセキュリティが最初から連携する包括的なネットワークフレームワークが不可欠です。ベライゾンは、APAC地域の企業向けにカスタマイズされたツールと専門知識を提供し、こうしたフレームワークの実現をサポートします。

ベライゾンがサポートする部門横断的なAIステアリングチームは、IT、セキュリティ、オペレーションの連携を図り、リスク評価や戦略の洗練に貢献します。Verizon Threat Research Advisory Center (VTRAC) は、リアルタイムの脅威インテリジェンスを提供し、攻撃の兆候となる異常なネットワークアクティビティなどを迅速に検出することができます。また、ベライゾンのサイバーセキュリティソリューションの一部である侵入テストなどのサービスは、脆弱性が悪用される前にギャップを特定し、生成AIプロジェクトの安全な立ち上げをサポートします。

AIツールが診断を支援する医療分野では、精度が非常に重要です。ロードアイランド州で行われたパイロットスタディでは、医師が不完全なAI出力に頼った場合、偽陽性率が86%に達したという報告がされました。ベライゾンのネットワークセキュリティサービスはデータの完全性を保護し、開発者が信頼性の高いAIシステムをゼロから構築できるようサポートします。

オーストラリア産業科学資源省は、システム設計上の欠陥やデータの問題がAIの出力に悪影響を及ぼす可能性があるという警告をしています。ベライゾンのサイバーセキュリティ評価は、これらのリスクを定量化し、プロセスを過度に複雑化させることなく生成AIの導入を強化するための実用的なインサイトを提供します。

AIセキュリティにおけるベライゾンの強み

ベライゾンはリスクを特定するだけでなく、その解決も支援します。ベライゾンのプライベート5Gおよびセキュアゲートウェイサービスにより、企業はAIワークロードのための安全な基盤を構築できます。また、Verizon Threat Research Advisory Centerは24時間365日体制の監視を提供し、お客様のビジネスチームが見落とす可能性のある脅威を組織が捕捉できるようにします。さらに、侵入テストとリスク評価は、APAC地域特有の課題に合わせてカスタマイズされた多層防御を実現します。

この文書の目的は、生成AIへの不安を煽ることではなく、備えの重要性をお伝えすることにあります。ベライゾンのソリューションは、開発初期段階からセキュリティを組み込み、イノベーションと保護のバランスをとることで、生成AIプロジェクトの成功をサポートします。

詳細はこちら

APAC全体にわたってエンタープライズAIのセキュリティ確保に貢献するベライゾンのサポートをご確認ください。

ベライゾンビジネスの営業担当者にメール（apaccontactus@verizon.com）で問い合わせるか、[verizon.com/business/en-au/contact-us/](https://www.verizon.com/business/en-au/contact-us/)にアクセスしてください。



