

生成AIを活用したリスクマネジメント

ソフトウェア品質保証プロフェッショナルの会 チーム6

上田 浩	(エンカレッジ・テクノロジー株式会社)
栗原 純一	(株式会社日立ソリューションズ・クリエイト)
菅原 広行	(ソニーセミコンダクタソリューションズ株式会社)
津久井 秀樹	(株式会社エクスマーシオン)
永田 哲	(元キヤノン株式会社)
深水 廉子	(SCSK株式会社) (発表者)
(五十音順)	

目次

1. はじめに
2. 組織の課題と本活動のゴール
3. 本活動のスコープ
4. 活動内容
5. 本活動のまとめ
6. 今後に向けて

1. はじめに

昨年度は「高品質なソフトウェアやシステムを開発・提供し続けられる組織」へと成長させるために、ノウハウの有効活用によって組織全体の技術力向上と過去の失敗の再発防止を目指し、生成AIの適用可能性を検証する活動を行った。

本年度は「**プロジェクトの開発ライフサイクルにおけるリスクマネジメント**」を主要テーマとし、**具体的な事例で生成AIの活用方法**を検証した。

本発表では、これらの活動結果を生成AIの活用事例として紹介する。

2. 組織の課題と本活動のゴール

■ 組織の課題

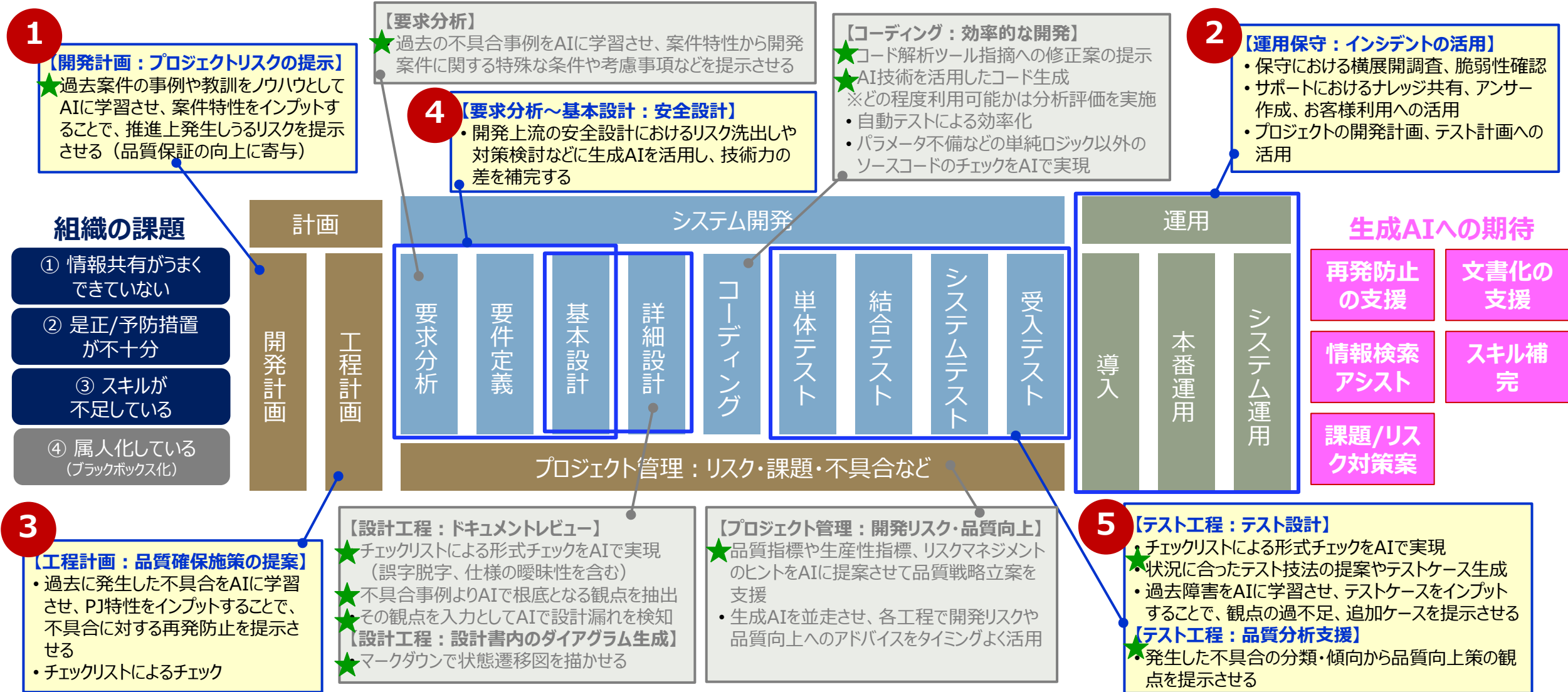
昨年度から取り組んでいる「先人の知恵（智慧）」や「過去のトラブル対応経験」といったノウハウを活かし、同じ失敗を繰り返さない組織を目指す中で、「**プロジェクト活動や開発ライフサイクルプロセスにおけるリスクマネジメント**」は、多くの組織に共通する大きな課題であると考えた。

■ 本活動のゴール

生成AIを活用し、社内外のノウハウや失敗事例を参考にすることで、**組織の誰もが必要かつ十分なレベルでリスクマネジメントを実施できるようになる。**

3. 本活動のスコープ

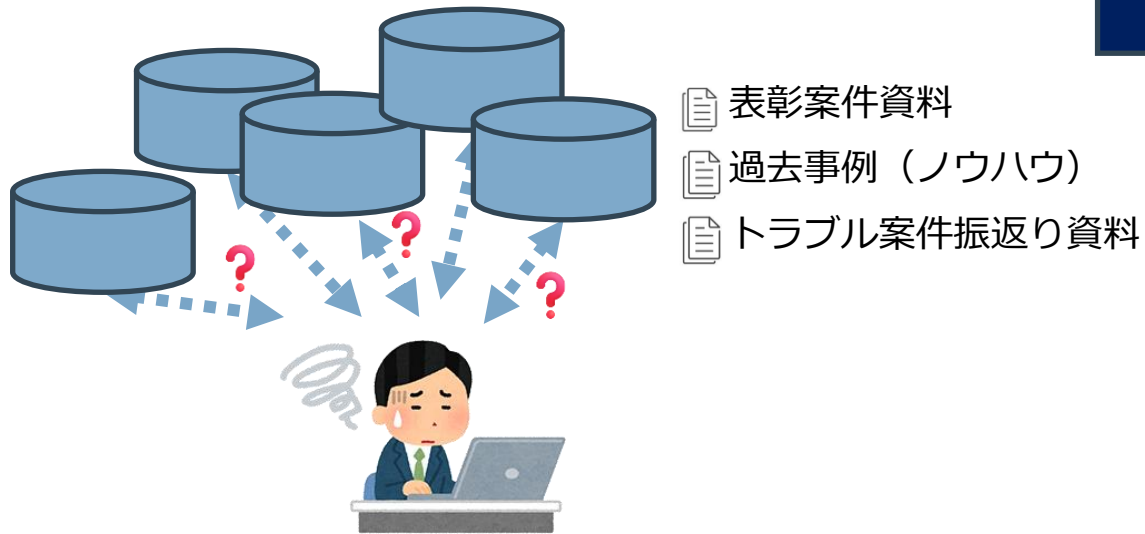
- 第16期活動の対象
- ★ 第15期活動の対象



4. 活動内容 ① 過去事例・教訓のAI活用によるリスク管理向上

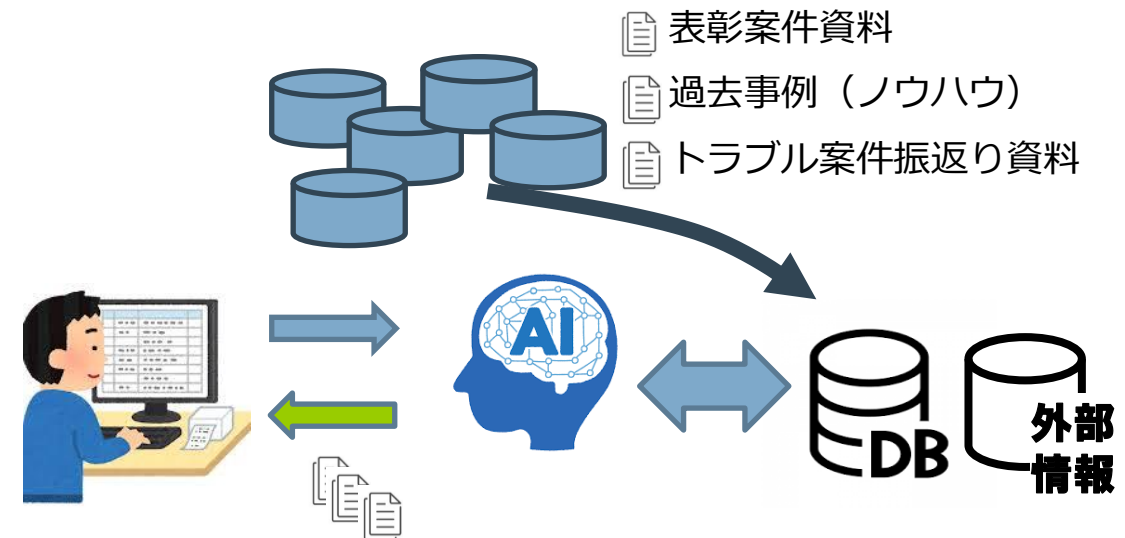
■ As-Is

- 社内ナレッジが散在しており、過去事例など知りたい情報が簡単に検索できない
- プロジェクト運営におけるリスク管理で過去のノウハウをうまく活用できない
- リスク管理は属人的で類似トラブルが再発している



■ To-Be

- 生成AI による過去事例を一元化し、一括検索することにより想定リスクを提示できる
- 関連事例を複数、参考文献のリンク先も表示することで詳細な過去事例の内容を確認できる
- システム開発関連の外部情報も併記し、多角的な情報を提供する



4. 活動内容

① 過去事例・教訓のAI活用によるリスク管理向上

■ 利用環境

Azure OpenAI (GPT-4o)

■ 現在の状況

- 検索ワードに適合する過去プロジェクトの事例を知ることができる
- キーワードや入力形式によりヒットする場合とヒットしない場合がある（検索結果のばらつき）
- β版として順次一般ユーザーに開放し、不具合の解消や改善要望の取り込みにつなげる予定

■ 課題と今後の対応

- 学習データの改善：マスキング不備、少ないデータ量、データ内容のばらつき
- RAG検索精度の向上：プロンプトの変更、AIモデルのバージョンアップ など
(GPT4o→GPT4.1or4.5,o3,o3pro)

4. 活動内容 ② インシデント情報を活用した生成AIによるリスクマネジメント

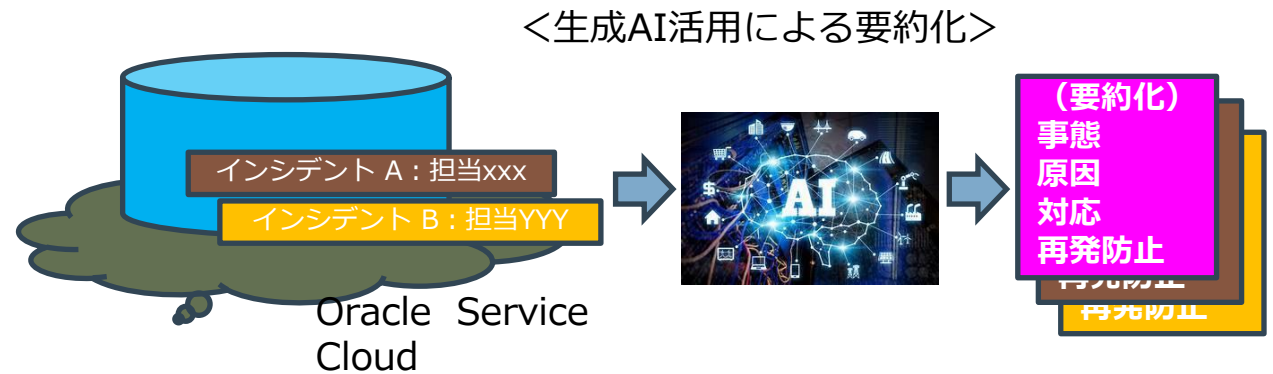
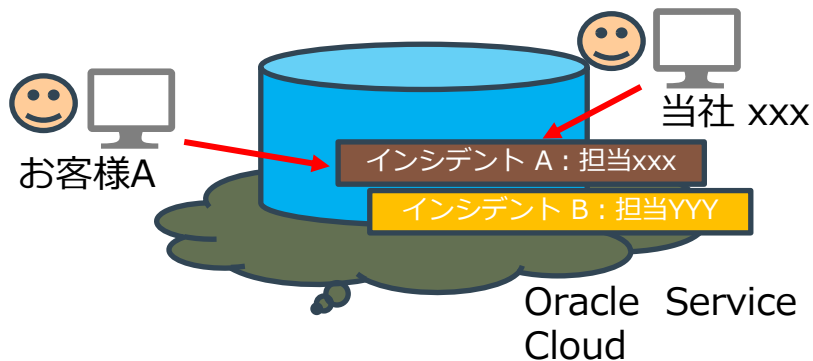
■ As-Is

- 自社開発のソフトウェア製品に対するインシデント情報が有効活用されておらず、対応策や再発防止策が個別最適の状況となっているため、全体品質の向上につなげていない。
- インシデント情報が各担当者の経験やスキルに応じて記載内容にバラツキがあり、データの再利用や情報共有化への支障となっている。



■ To-Be

- 生成AIを利用したインシデント情報の要約を行い、ナレッジの共有化とデータ利用時の信頼性向上を促進する。さらにFAQ情報、マニュアルを取り込んだRAG+AI利用による開発計画・テスト計画策定やプロジェクト管理におけるcopilot利用を行い、全体品質向上へ寄与する。



4. 活動内容 ② インシデント情報を活用した生成AIによるリスクマネジメント

■ 利用環境

Azure OpenAI (GPT-4o)、Oracle Service Cloud (インシデント情報基盤)

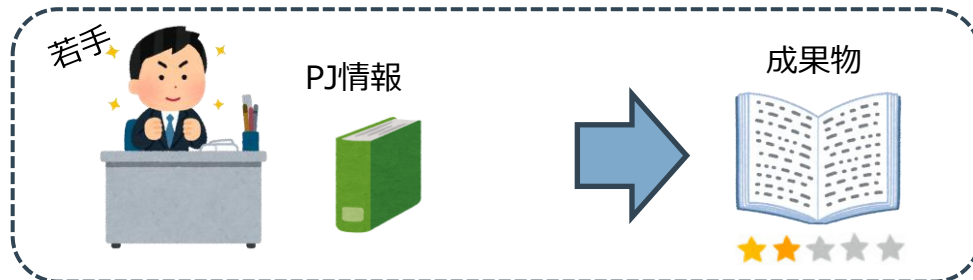
■ 現在の取り組み状況と今後の予定



4. 活動内容 ③ AIによる品質保証計画の策定補助

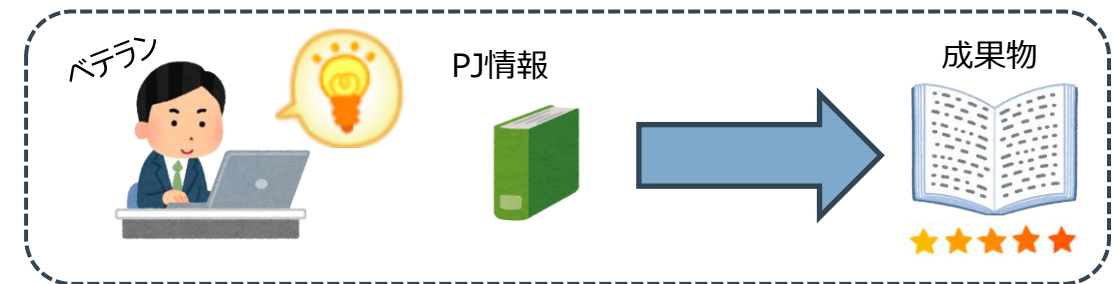
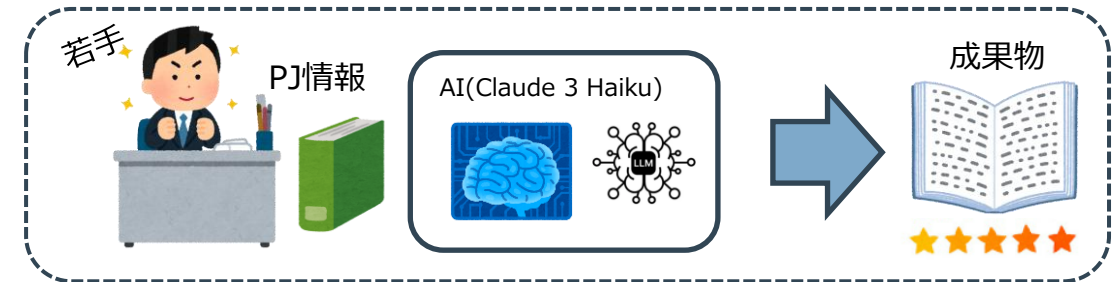
■ As-Is

- QA人財不足が顕著であり、生成AIを活用した作業の効率化が急務。
- 経験の浅い若手QAとベテランQAとで作業成果に差が大きい。



■ To-Be

- 生成AI による若手QAの作業支援
- 若手QAに不足している経験を生成AIが補足。
- ベテランQAと近いレベルの成果物が作成可能。



4. 活動内容

③ AIによる品質保証計画の策定補助

■ 利用環境

社内生成AI（ Claude 3 Haiku ）

■ 現在の状況

方針は以下の通りとして、対応を進める。

- 👉 対象の成果物 : 品質保証計画書
- 👉 生成AIによる支援 : QA経験が大きく影響する「品質保証の勘所」を作成

■ 今後の対応

複数のプロジェクトで若手QAに「勘所」を作成させ、ベテランQA作成の「勘所」の差異を確認する。

確認結果に従い、今後の対応方針を定める。

⇒ 差異「有」 : 差異を埋めるための生成AI活用方法を再考する。

⇒ 差異「無」 : 「勘所」を用いて、品質保証計画書作成が可能か確認を行う。

4. 活動内容 ④ 開発上流での安全設計への生成AIの活用

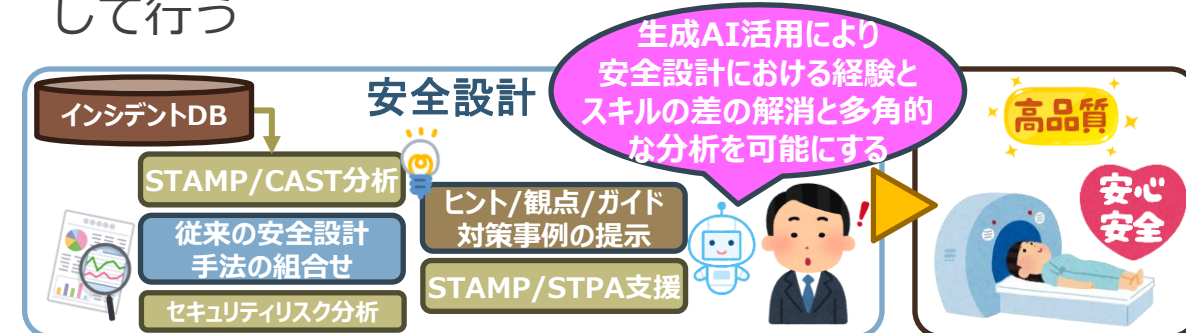
■ As-Is

- 近年のシステムは高度化・複雑化が進み、医用機器などの安全設計において、FMEA、FTA、HAZOPなどの従来手法だけでは、潜在的な新たなリスクを十分に特定できない可能性がある
- システムの複数の構成要素間の相互作用によって起こるアクシデントを分析するSTAMP/STPAについても、実用例が限られている
- 最近では、セキュリティの脆弱性が原因の安全性に影響を及ぼすアクシデント事例も報告されている
- 現状のやり方では、個人のスキルや経験に依存する部分が大きく、分析の質や網羅性にばらつきが生じ、重要なリスクが見落とされる



■ To-Be

- リスク抽出の観点や、検討漏れの発見・気づき、制約事項、類似インシデント事例などの提示
- 対策検討時の支援やリスクコントロール・対策事例などの提示、設計フレームワークなどの提供
- ミスユースケースの抽出、セキュリティ脆弱性が原因で発生する安全のリスクガイド (医療機器の例：システムの暴走、X線の過剰被曝、停止など)
- リスクの重大性(問題発生時の影響度)の基準例
- 中・長期的な未然防止・再発防止のための手段・対策案検討のアドバイスや類似案件の検索提示
- 製品のインシデントなどの社内の機密情報の参照
- 従来手法で構成要素の故障や事象からの分析も並行して行う

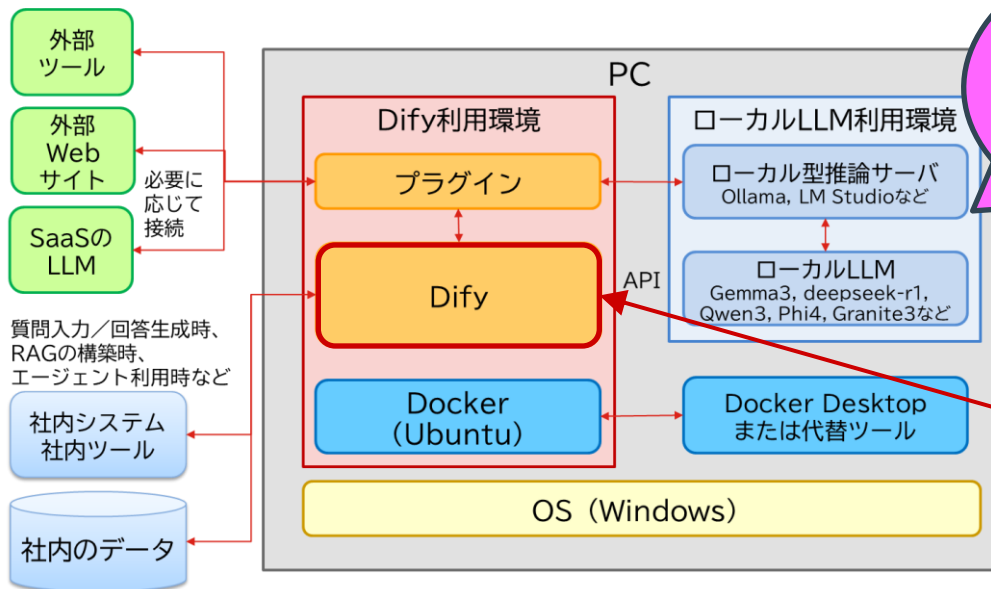


4. 活動内容 ④ 開発上流での安全設計への生成AIの活用

ローカルLLM評価環境で検証

仮説：STPAでの経験や技術力の差が解消できる

- Windows 11のNote PCにローカル環境を構築
- 約1時間でセットアップ可能（アプリ作成時間除く）
- LLMはGemma3ほか、複数の4~8bのモデルを利用
- Difyのチャットフローでローカル環境を構築し、複数のLLMをパラレルで利用して回答結果を比較検証したり、RAGで回答精度の向上を図る



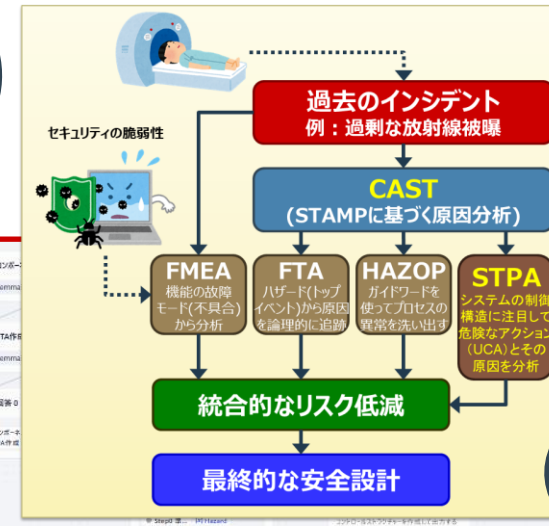
評価環境はNote PCに簡単にセットアップ
LLMやアプリはすべてフリーソフトを利用してノーコードで構築



安全設計への適用結果

結果：期待通りの効果 - 仮説が検証できた

- アプリ検証：Difyのチャットフローで下記を確認
 - ① CASTで事故とセキュリティ脆弱性を分析 → RAGに登録
 - ② ハザードと安全制約を定義（セキュリティ含む）
 - ③ 制御構造と因果関係を分析
 - ④ FMEA・FTA・HAZOP・STPAと連携して対策案を導出
 - ⑤ セキュリティ設計を含めて安全設計に反映



チャットフローの中のRAGとIn-Context Learningを使って社内ノウハウを取り込んで利用

RAGやチャットフローの更なる改善など一部課題は残るものの効果が確認できた

4. 活動内容

⑤ テスト工程での生成AI活用とセキュリティスキル向上

■ As-Is

- ソフトウェア製品の開発・テストメンバーのセキュリティスキルが低い
- セキュリティの専門家をタイムリーにプロジェクトに参加させるのはコストとタイミング的に困難
- その結果、セキュリティ品質が低い製品・サービスをリリース



■ To-Be

- 生成AIを活用して、レビュー、単体テスト、ファジングテスト、E2Eテストを作成しながら、セキュリティのスキルアップをプロジェクトメンバー全員が行う
- 全員がOWASP Top10などのセキュリティの基礎知識を取得後、生成AIにWSTGをベースにセキュリティ対策とその確認テストを計画・設計させる
- 生成AIの出力を鵜呑みにせず、不明点、漏れを問い、アプリに対してテスト実行する
- なぜテストがOKで終わったかを自分で考え、AIの説明が納得できるまで更にテストを生成させたり、必要であればアプリを変更させたりする

4. 活動内容 ⑤ テスト工程での生成AI活用とセキュリティスキル向上

■ 生成AIの自信満々の提案を鵜呑みにしてはいけない

- コード生成AIでは“沼にはまる”ことがしばしばだが、絶対に「できない」とは言わずに、自分で書いたコードをどんどん削除したり、日本語のコメントを英語にしたり、try-catchで囲んでcatch句に余計なコードを書いたりする。それを見つけて

人「実質何もやっていないのと同じだ。」

AI「申し訳ありません。」と謝るが最後に

AI「直らないのはユーザーの環境が悪い。」

と言い出すこともあるので、そんな時は自分で問題解決して、AIに人間様の力をみせつける

- トークンをどんどん消費するように作られている。私のClaude 3.5, Gemini2.5Proの経験、「そろそろこいつの限界」を見極めることが重要。

- また格納型XSSのテストコードの生成でも

AI「意図的に脆弱なコードにしてもXSSを確認できないのは、アプリのどこかにまだ防御コードがあるに違いありません。」

人「いや、テストコードのログには発生を検知したとあるので、アプリでなくテストコードの問題で発生確認ができないのでは？」

AI「了解。XSS発生のダイアログを即消しているのでも人間の目で見えないようです。消すのをやめますので確認してください。」

人「ちゃんと確認できました。」

■ 利用環境：Claude 3.5, Gemini2.5Pro など

5. 本活動のまとめ

本活動を通して、下記のように様々な局面で、生成AI活用の有効性を確認することができた。

- 多種多様なフォーマットの事例データも、AIに取り込み一元化するだけでも、検索効率の向上、ノウハウの共有といった効果はある
- 膨大なインシデント活用には、生成AIの要約でデータを均質化するだけでも十分効果が見込め、将来的にはリスクマネジメントへの活用も期待できる
- 生成AIにより経験不足を補完し、業務活動の質の向上が期待できる
- 安全設計において、チャットフロー＋RAGで、社内の機密情報やノウハウが活用できる目的特化型の生成AI利用環境を提供することにより、スキルと経験の差を解消できる
- LLMの活用で、セキュリティテストコードの自動生成と担当者のスキル向上を同時に達成できる

6. 今後に向けて

■ 本活動の振り返り

- 各社**取り組み状況を共有**しながら進められたことで、**自社での展開・方向性**に活かした
- 過去ノウハウのナレッジベースをつなげて、適用するリスク分析手法に沿って生成AIに役割・プロンプトを与えることで、**リスク分析の伴走者**としての活用に有用性を感じた
- **生成AIを使うことを目的とするのではなく**、導入効果が出せる戦略を立て、仕組みや運用プロセスを構築した

■ 今後に向けて

- 手法に応じた**プロンプトの標準化**
- **AIエージェントによる意思決定の自動化**
- **単体テストとセキュリティ用ファジングテストの自動生成** や **セキュリティ設計とその実装レビューへのLLM活用**
- **品質保証活動の改善**に向けて、**生成AI**などの新技術の積極的な取り込み
- **経営企画におけるリスク洗出し**や**プロダクトのテストへの応用**

ご清聴ありがとうございました