

安全保障技術研究推進制度 令和6年度終了課題 終了評価結果

1. 評価対象研究課題

- (1) 研究課題名：強化学習を用いた環境適応型ファジングシステムの提案
- (2) 研究代表者：株式会社リチエルカセキュリティ 木村 廉
- (3) 研究期間：令和2年度～令和6年度

2. 終了評価の実施概要

日時：令和7年12月12日

場所：TKP新橋カンファレンスセンター

評価委員：未来工学研究所 理事長、上席研究員

東京大学 名誉教授

平澤 洽（委員長）

リモート・センシング技術センター 理事

岩野 和生

玉川大学 名誉教授

大森 隆司

東京工業大学（現 東京科学大学） 名誉教授

佐藤 誠

兵庫県立大学大学院 情報科学研究科 教授

田中 俊昭

千葉商科大学 特定教授

筑波大学 名誉教授

東京工業大学（現 東京科学大学） 名誉教授

寺野 隆雄

札幌市立大学 理事長、学長

中島 秀之

山口大学 大学院創成科学研究科 教授

西井 淳

産業技術総合研究所 人間社会拡張研究部門 上級主任研究員

長谷川 良平

福井工業大学 経営情報学部長・主任教授

大阪大学 名誉教授・特任教授

馬場口 登

（委員長以外は五十音順・敬称略）

3. 研究と成果の概要

研究の概要

ソフトウェアの脆弱性が社会基盤や国家安全保障に深刻な影響を及ぼす現代において、脆弱性を迅速かつ効率的に検知・修正する技術の確立は喫緊の課題である。特に、修正プログラムが存在しないまま攻撃者に悪用される「ゼロデイ脆弱性」の早期発見は、被害の拡大を防ぐ上で極めて重要である。

本研究では、ファジング技術に着目し、アプリケーションや IoT 機器といった多種多様なプログラムからゼロデイ脆弱性を発見する基盤の構築を目指した。

ファジングはこれまで、(1)最適なアルゴリズム選択の難しさ、(2)クラッシュの原因特定や脅威度評価手法の不足、(3)ブラックボックス環境における適用困難、(4)ハードウェア支援機構の活用限界といった課題を抱えてきた。本研究では、これらの課題に対し、強化学習の適用可能性の解明を軸として、理論的検討と実用的手法の両面から包括的に解決を図った。

成果の概要

本研究では、以下の主要な成果を達成した。

- ・**統合ファジング基盤「fuzzuf」の開発**: 13 種類の主要ファジングアルゴリズム (AFL、AFLFast、VUzzer、libFuzzer、IJON、Nezha、DIE、Nautilus、MOpt、SymCC、Eclipser、K-Scheduler、SLOPT) を統合し、異なる手法間での比較検証や組み合わせを可能にする環境を実現した。また、ファジングアルゴリズム記述用 DSL 「HierarFlow」を新たに提案し、柔軟なプロトタイピングを可能にした。
- ・**強化学習の適用可能性解明と最適化手法の確立**: ミクロ最適化手法「SLOPT」およびその発展型「rezzuf」を提案し、Google の OSS-Fuzz が長期間検証しても未発見だったゼロデイ脆弱性を含む計 26 件の新規脆弱性を発見した（うち 17 件が CVE 登録、5 件がクリティカル）。
- ・**クラッシュ原因特定・脅威度評価基盤の構築**: 既存手法の比較検証プラットフォーム「RCABench」を開発し、NDSS 併設 BAR ワークショップで Best Paper Award を受賞した。また、「任意コード実行 (RCE) 可能性」に基づく脅威度評価指標を確立し、57 件の既知脆弱性で人手分析と同等以上の精度を達成した。
- ・**ブラックボックス環境向けカバレッジ推定技術**: 静的解析結果と通信レスポンスを組み合わせた手法「Shepherd」を開発し、既存手法比で最大 54.8% のカバレッジ推定精度向上を達成した。ISSTA 併設 Fuzzing Workshop に採択された。
- ・**重要インフラへのファジング適用可能性の体系化**: NISC が定める重要インフラ 15 分野を対象に、ファジングの技術的適用可能性と社会的・制度的制約を横断的に整理した初の体系的調査を実施した。
- ・**人材育成基盤の構築**: ファジングトレーニングプログラムを 2 社 17 名に提供し、SECCON 13 でのワークショップでは 34 名が参加、最優秀賞（オールラウンドプレ

イヤー賞) を獲得した。

4. 終了評価の評点

A A 想定以上の成果をあげた。

5. 総合コメント

ファジングを高性能で行う新たなフレームワークを提案し、計画以上の成果を得るとともに、ファジングの限界や脆弱性の新たな評価手法を示す等、高いレベルで目標を達成したことは評価できる。新しいアルゴリズムの開発による当初の課題の多くに成果を出しており、採択の成果があったものとする。

生成 AI によるファジングの検討も盛んになりつつあるので、今後、より発展的な研究を進めることを期待する。また、特定の OT (Operational Technology) の分野などで具体的なインパクトを出すための構想や、戦略的なアライアンスを通じて次のステップを明確にし、インパクトを起こせる形に発展させていきたい。

さらには、外部との連携を増やすことで、活動が広がることを期待する。本技術を核にしたビジネスモデルについても検討いただきたい。

6. 評価の観点ごとの評価結果と個々の委員によるコメント

6-1. 研究開始時に設定した研究目標の達成度 (主題的成果)

(1) ファジングフレームワークの設計、開発【最終目標を上回る成果】

当初目標であった 8 種類のファジングアルゴリズムの統合を大きく上回り、13 種類のアルゴリズムを統合した汎用プラットフォーム「fuzzuf」を開発した。さらに、既存手法の構成要素を柔軟に再利用・再構成するための独自記述言語「HierarFlow」を新たに提案し、多様な構成の迅速なプロトタイピングを実現した。

fuzzuf と HierarFlow を基盤として、SLOPT と K-Scheduler を組み合わせた新たなファジングアルゴリズム「rezzuf」を提案し、1 週間の試行で 23 件のゼロデイ脆弱性を発見した。この成果は国際セキュリティカンファレンス HITCON にて発表した。

なお、「複数のファジングアルゴリズムを統合可能とするフレームワーク」という発想は、同時期に国際会議 CCS で提案された LibAFL にも見られ、本研究が同時代の最有力研究グループと同等の見地に独自に到達していたことを示している。

(2) 強化学習エンジンの開発、検証【最終目標に相当する成果】

ファジングへの強化学習適用について、「ミクロな最適化」(変異操作の最適適用順序・頻度の探索) と「マクロな最適化」(複数アルゴリズムの選択最適化) の両面から検証を行った。

ミクロ最適化では、複数の変異操作を組み合わせた結果が適用順序に依存しない

という重要な性質を発見し、探索空間の大幅な削減に成功した。これにより提案した「SLOPT」は、デファクトスタンダードである AFL++ を超える探索効率を示し、Google の OSS-Fuzz が長期間検証していたソフトウェアから新規ゼロデイ脆弱性 3 件を発見した。SLOPT 論文は国際会議 ACSAC（採択率 24%、Core Conference Rank : A）にて発表した。

マクロ最適化では、強化学習エージェントによる動的なアルゴリズム選択手法を提案し、類似プログラム間での転移学習可能性を実証した。一方で、「強化学習にかける時間をファジングに使う方が実用的」という現実的知見も得られた。

これらの成果により、従来「未解決問題（オープンプロブレム）」とされていた「ファジングにおける強化学習の実用化」は「解かれた課題」となった。

(3) アクセラレータの開発、検証【最終目標を上回る成果】

ARM CoreSight を活用した高速化手法を構築し、ソースコードがない場合において約 1.6 倍の高速化を実現した。ただし、ソースコードがある場合は Intel 環境と比較して効果が限定的であり、その原因が Intel と ARM の CPU におけるレジスタ呼び出し仕様の違いという根本的な制約にあることを解明した。これはハードウェア設計上の限界であり、ARM 環境向けアクセラレータのこれ以上の高速化は原理上困難であるという重要な知見である。

RISC-V 環境については、(1) コロナ禍に起因する半導体不足、(2) プロセッサトレース仕様が議論中で後に大きく変更される可能性、(3) 実世界の IoT 機器での採用事例の少なさから、計画を軌道修正した。

この反省を踏まえ、CPU 機能への依存を前提としないブラックボックスファジングの効率化にアプローチを転換し、副次的成果の創出を試みた。

【個々の委員によるコメント】（主題的成果）

- ・組み込みデバイスに対する優位性の検証は研究開始時の設定した目標件数（15 件）に対して十分な検討はなされていないが、ファジングの最適化システムの構築に関しては当初目標を十分達成している。
- ・現実的なサイバー犯罪をモデル化したファジング状況における脆弱性の発見技術の開発で大きな成果を挙げている。
- ・目標は明確で十分な成果が出ていると評価できる。

6-2. 計画時に想定していなかった成果（副次的成果）

(1) クラッシュ原因特定・脅威度解析技術の確立

研究進行中に、ファジングで得られる大量のクラッシュのうち実際に攻撃可能な脆弱性に結びつく事例は限られるという認識に基づき、クラッシュ原因・脅威度解析を新たに研究課題として設定した。

クラッシュ原因特定については、Aurora および VulnLoc といった既存手法を統合した比較検証プラットフォーム「RCABench」を開発した。RCABench 論文は NDSS 併設 BAR ワークショップにて採択され、Best Paper Award を受賞した。RCABench を通じて、既存手法が近傍の異なる原因箇所を同一と誤識別する問題や、変数値自体が原因のバグの根本原因を特定できない問題を解明した。

脅威度評価については、「任意の場所に任意の値を書き込めるクラッシュほど脅威度が高い」という実用的かつ直感的な評価指標を導出した。57 件の既知脆弱性への適用により、すべての脆弱性について人手評価と同等以上の精度を達成した。さらに、人間による脆弱性脅威スコア (CVSS=6.5) が中程度であっても任意コード実行が可能な事例を特定し、専門家評価を超える分析能力を実証した。

(2) IoT 機器のブラックボックスファジング効率化

ブラックボックスファジングではカバレッジが得られずグレーボックスファジングほどの性能が発揮できないことを踏まえ、代替指標を検討し、レスポンス情報と静的解析結果を組み合わせた擬似カバレッジ推定手法「Shepherd」を提案した。Shepherd は既存手法 Labrador と比較して推定精度を平均 30.2%、最大 54.8%向上させ、無線 LAN 中継器から 2 件のゼロデイ脆弱性を発見した (1 件に CVE 登録)。論文は ISSTA 併設 Fuzzing Workshop に採択された。

ARM・RISC-V 双方の環境で CPU 機能を活用したアクセラレータ実現の限界を科学的に分析し、その知見をもとに IoT 機器に対する真に実用的なブラックボックスファジング技術を確立したことは、当初の最終目標を質的にも量的にも上回る成果である。

(3) 重要インフラへのファジング適用可能性の体系化

NISC が定める重要インフラ 15 分野 (情報通信、金融、航空、空港、鉄道、電力、ガス、政府・行政サービス、医療、水道、物流、化学、クレジット、石油、港湾) を対象に、ファジングの技術的適用可能性と社会的・制度的制約を横断的に整理した。

調査の結果、情報通信・電力・ガス・水道など、プロトコル仕様や OSS 実装が豊富な分野では継続的適用が比較的容易である一方、鉄道・航空・クレジットといった閉鎖性の高い分野では実機への適用が困難であることが明らかとなった。また、上水道施設の PLC を対象としたファジング検証を通じて、EDSA 認証要件との接続可能性を実証した。

我々の知る限り、本調査は重要インフラ 15 分野すべてを対象とした初の体系的試みであり、分野ごとの特性を踏まえた構造的な意思決定を可能にする基盤を提供した。

【個々の委員によるコメント】（副次的成果）

- ・計画時と状況は大幅に変化している。その変化への対応は必ずしも十分とは思えない。
- ・本手法の適用範囲と限界、クラッシュ原因分析・脅威度の自動化など、成果が出ている。
- ・脆弱性発見の訓練と専門家の育成の努力については評価できる。
- ・脅威度判定の自動化やブラックボックスファジング効率化など多数の副次的成果を得ている。
- ・クラッシュ原因などの解析も可能になったのは、セキュリティシステムの開発者にとっても有用な情報となりうる。

6-3. 他の者により派生した成果（間接的成果）

本研究のファジングトレーニングプログラム受講者から、国際セキュリティカンファレンスへの研究採択を達成した事例が報告された。これは、本研究の教育的波及効果が他の研究グループにまで及んでいることを示している。

また、本研究で構築した知見や技術基盤は、NICT、産業技術総合研究所、筑波大学、民間セキュリティ企業等の研究者・技術者との意見交換を通じて共有され、各機関の研究活動に間接的に貢献した。

【個々の委員によるコメント】（間接的成果）

- ・社外展開としての企業向けトレーニングの開催や、ハッキングコンテストにおけるワークショップ開催は評価できる。
- ・NICT や AIST などとの積極的な情報交換も行っていることで全体的な成果の底上げにつながっている。

6-4. 科学技術上特筆すべき成果

(1) ファジング分野における複数のオープンプロブレムの解決

本研究は、従来個別に扱われてきた複数のファジングアルゴリズムを統合可能な実行基盤を構築し、強化学習の適用可能性を理論・実証の両面から解明し、ハードウェア支援機構の限界と代替手法を開発することで、ファジング分野における複数の未解決課題を同時に解決した。特に、「ファジングにおける強化学習の実用化」というオープンプロブレムを「解かれた課題」に転換したことは、学術的に特筆すべき成果である。

(2) 実世界での脆弱性発見実績

提案手法群により合計 26 件のゼロデイ脆弱性を発見し、うち 17 件が CVE として登録、5 件がクリティカル（深刻度高）と評価された。特に、Google の OSS-Fuzz が

長期間検証していたソフトウェアから SLOPT のみが検出可能な脆弱性を発見したことは、提案手法の独自性と有用性を実証するものである。

(3) 競争的研究環境での先見性

本研究期間中、ファジング分野では大きな技術革新が進行した。2024 年には、米 DARPA の「AI Cyber Challenge (AIXCC)」において、AI による自律的な脆弱性発見・パッチ生成技術が実証され、XBOW のような AI 駆動型ペネトレーションテスターが HackerOne で 1 位を達成した。Google の Big Sleep も、従来のファジングでは発見できなかった SQLite のゼロデイ脆弱性を AI で発見した。

このような急速な技術進歩の中で、本研究が先行して確立した「複数アルゴリズムの統合基盤」「強化学習の適用」「脅威度の自動評価」という技術基盤は、今後の AI × ファジング融合研究の土台として重要な位置を占める。本研究は、これらの技術トレンドを先取りする形で成果を挙げており、国際的な競争環境における先見性を示している。

【個々の委員によるコメント】(科学技術上特筆すべき成果)

- ・ファジングシステムの適用範囲を大きく進展させた。
- ・ファジングのアルゴリズムの改善で、競合の他アルゴリズムを凌駕することを示した。
- ・ファジングの適用範囲と限界について、学会での複数件の受賞につながった。

6-5. 論文（投稿中のものも含む）、学会発表等

雑誌論文：4 件

口頭発表：6 件

展示：2 件

図書（技術解説記事）：14 件

プレス：1 件

その他（脆弱性報告）：26 件

6-6. 特許（出願中のものも含む）

出願中：1 件

【個々の委員によるコメント】(論文、特許、学会発表等の研究成果)

- ・Conference 論文、学会発表など、外部発表が行われており、賞をいただくなど評価されている。一方で、技術報告がブログ記事であり、普及範囲が狭い。
- ・ACSAC や IEEE S&P など著名な国際会議で発表されている。
- ・雑誌論文として報告されているものはいずれも学会の Proceeding であり、タイプ

Sの成果としてはやや少ない印象を受ける。

6-7. 科学技術への波及効果

(1) 後継プロジェクトへの展開

本研究の成果は、NEDO 委託事業「経済安全保障重要技術育成プログラム／先進的サイバー防御機能・分析能力強化」(JPNP24003)において、以下の2つの研究テーマとして発展的に継続されている。

- ・ **IoT 機器に対するトリアージ技術**：本研究で開発した Shepherd を発展させ、ブラックボックス環境でのクラッシュ収集・入力最小化・脅威度判定技術の確立を目指す。
- ・ **脆弱性探査の阻害技術**：本研究で蓄積したファジング阻害要因の知見を活用し、アンチファジング技術を含む耐タンパー性獲得技術の研究開発を行う。

※一般社団法人サイバーリサーチコンソーシアムからの再委託先／研究分担組織として

(2) 社会実装への展開

本研究の成果は、以下の形で社会実装が進められている。

- ・ **セキュリティサービスへの統合**：開発したファジング技術をバイナリ解析や IoT 機器診断サービスに統合し、従来の手動診断では発見困難だった脆弱性の検出や効率的なトリアージを実現する。
- ・ **人材育成プログラムの継続提供**：ファジングトレーニングプログラムを継続的に提供し、国内のファジング人材育成基盤を拡充する。
- ・ **コンサルティングへの知見活用**：重要インフラ 15 分野へのファジング適用可能性に関する知見を、各分野の特性に応じたセキュリティコンサルティングに転用する。

(3) 学術的波及効果

本研究で確立した技術基盤は、今後のファジング研究における標準的な比較・評価基盤として活用されることが期待される。特に、fuzzuf は複数のアルゴリズムを公平に比較可能な環境を提供し、RCABench はクラッシュ原因特定手法の評価基準を確立した。これらは、ファジング分野の研究の質向上と再現性確保に貢献するものである。

【個々の委員によるコメント】(発展性と波及効果)

- ・ 技術力は評価されたが、どのように継続性をもって世の中に影響を与えていくかのモデルがまだ判然としていない。
- ・ 人材の育成による新たな成果が発生しており、また次のプロジェクトに採択されるなど、波及効果が出ている。

- ・この分野の人材育成にも役立っている点は評価したい。
- ・開発したフレームワークを公開しており、セキュリティ監査サービスへの展開が期待される。
- ・サイバー犯罪の取り調べにも貢献可能と思われる。
- ・セキュリティ分野の実ビジネスにつなげていただきたい。

6-8. 効率的な研究実施体制とマネジメント

本研究では、継続的な進捗把握と柔軟な課題修正を可能とする体制構築に注力し、成果創出を支える運営体制を確立した。

- ・**週次研究会議**：原則として週次で進捗報告・課題抽出・アクション設計を行い、メンバー間での認識統一を図った。
- ・**ツール活用による効率化**：HackMD・Trello からプロジェクト複雑化に伴い Notion へ移行し、情報の集約と可視性を向上させた。リモートワーク環境下では Gather および Google Meet を活用し、分散拠点間での即時意見交換を実現した。
- ・**間接業務の削減**：Slack 上のワークフロー整備により、経費関連の申請・承認・管理プロセスを簡素化し、研究者の本来業務への集中を可能にした。
- ・**機動的なチーム編成**：研究の進展状況や技術課題に応じてプロジェクトメンバーの構成を再編し、必要に応じて集中作業のための合宿も実施した。
- ・**外部有識者との連携**：NICT、産業技術総合研究所、筑波大学、民間セキュリティ企業等の研究者・技術者との意見交換を行い、方針の妥当性を検証した。

【個々の委員によるコメント】（効率的な研究実施体制とマネジメント）

- ・発足時に対して大幅にメンバーの補強が行われ、それが挑戦的成果に結び付いた。
- ・当初のスタートアップの状態から社員 20 名にまで組織を拡大し、必要な研究を遂行していると同時に、ある程度の売り上げもあり、適切なマネジメントがなされている。
- ・少数の優秀なエンジニアに依存した体制であるため、その人材が流出するとその後どうなるのか懸念される。
- ・スタートアップの規模拡大を本研究資金をベースに実現している点は評価できる。

6-9. 研究推進時に生じた問題への対応

(1) RISC-V アクセラレータ開発の軌道修正

当初計画していた RISC-V 環境向けアクセラレータの開発については、以下の 3 つの理由から軌道修正を行った。

- ・コロナ禍に起因する半導体不足（FPGA ボードの納品予定が発注後 45 週間後）
- ・RISC-V プロセッサトレース仕様が議論中で、後に大きく変更される可能性
- ・実世界の IoT 機器での採用事例が少ないこと（当時）

この課題認識を踏まえ、CPU 機能への依存を前提としないブラックボックスファジングの効率化に研究方向を転換し、結果として Shepherd という成果を得ることができた。

(2)ARM 環境でのハードウェア支援限界の発見と対応

ARM CoreSight を活用したアクセラレータ開発において、ソースコードがある場合の効果が Intel 環境と比較して限定的であることが判明した。原因を精緻に分析した結果、Intel と ARM の CPU におけるレジスタ呼び出し仕様の違いが根本的な制約であることを解明した。この知見は、今後の実装戦略において CPU 支援型ファジングの現実的適用範囲を明確化し、代替手法の必要性を示す重要な成果となった。

【個々の委員によるコメント】（研究推進時に生じた問題への対応）

- ・この分野は次々と新たな課題が生み出されてきている。それに対応できる新たな発想を持った人材を確保すべきである。
- ・当初に想定した計画が適切でないとの判断から、副次的な問題に課題をシフトし、新たな課題で成果を出している。
- ・半導体不足や研究推進のための想定以上の工数等の問題に適切に対応して、適切な成果を出している。

6-10. 経費の効率的な執行

本研究では、必要な機材やリソースを適切に投入しつつ、外部環境の変化にも柔軟に対応しながら、効率的に経費を執行した。

- ・ **計算基盤の整備**：パーソナルスーパーコンピュータ DGX A100 を導入し、強化学習や長期間ファジングなどの計算集約型作業を効率化した。オフィスとデータセンター間の専用ネットワークを敷設し、解析結果の即時同期や大容量データ転送を可能にした。
- ・ **外部環境変化への対応**：COVID-19 や半導体供給難の影響で一部機材購入を見送ったが、計算リソース全体としては当初計画を満たす構成を確保し、研究遂行に支障はなかった。
- ・ **間接経費の有効活用**：進捗管理ツール、検証用機材、学生の研究補助員の動員に間接経費を活用し、開発サイクルの短縮と検証工数の削減を実現した。

【個々の委員によるコメント】（経費の効率的な執行）

- ・概ね適切と判断する。

【参考】用語集

<基本概念>

- ・ファジング (Fuzzing) : ランダムに生成した入力を大量にソフトウェアへ与え、クラッシュや異常動作を観察することで脆弱性を発見する動的テスト手法
- ・ゼロデイ脆弱性 : 修正プログラム (パッチ) が存在しない状態で発見・悪用される脆弱性。攻撃者に先に発見されると被害が拡大しやすい
- ・クラッシュ : プログラムの異常終了。ファジングではクラッシュを引き起こす入力を収集し、脆弱性の有無を調査する
- ・カバレッジ : プログラムの実行時に通過したコード領域の割合を示す指標。ファジングでは高いカバレッジを達成することで未発見のバグを探索する
- ・トリアージ : 大量のクラッシュから脆弱性の深刻度を判断し、対応優先順位を決定するプロセス

<ファジングの分類>

- ・グレイボックスファジング : プログラムの内部情報 (カバレッジ等) を一部利用しながら行うファジング。AFL 等が代表例
- ・ブラックボックスファジング : プログラムの内部情報を使用せず、外部からの入出力のみを観測して行うファジング。IoT 機器等に適用される
- ・カバレッジガイドファジング : 実行時のカバレッジ情報を活用し、新たなコード領域に到達する入力を優先的に探索する手法

<脆弱性・セキュリティ用語>

- ・CVSS (Common Vulnerability Scoring System) : 脆弱性の深刻度を 0.0~10.0 の数値で表す国際的な評価基準
- ・RCE (Remote Code Execution) : 攻撃者がリモートから任意のコードを実行できる脆弱性。最も深刻な脆弱性の一つ
- ・アンチファジング : ファジングによる脆弱性発見を妨害・困難にする技術

<本研究で開発した技術>

- ・fuzzuf : 本研究で開発したファジングフレームワーク。13 種類のアлゴリズムを統合し、比較・組み合わせを可能にする
- ・HierarFlow : fuzzuf 上でファジングアルゴリズムを記述するためのドメイン固有言語 (DSL)
- ・SLOPT : 強化学習の単純化モデル (バンディットアルゴリズム) を用いた変異操作の最適化手法
- ・rezzuf : SLOPT と K-Scheduler を組み合わせた発展型ファジングアルゴリズム
- ・RCABench : クラッシュ原因特定手法 (Root Cause Analysis) の比較検証プラットフォーム
- ・Shepherd : レスポンス情報と静的解析結果を組み合わせた、ブラックボックス環境向けカバレッジ推定手法

<既存ツール・基盤>

- ・AFL (American Fuzzy Lop) : カバレッジガイドファジングの代表的なツール
- ・AFL++ : AFL の発展版。コミュニティにより多くの機能が追加されたデファクトスタンダード
- ・OSS-Fuzz : Google が運営するオープンソースソフトウェア向けの継続的ファジング基盤。2023 年時点で 3.6 万件以上のバグを検出
- ・LibAFL : 本研究と同時期に CCS で発表された、複数アルゴリズムを統合可能なファジングフレームワーク

<ハードウェア関連>

- ・ARM CoreSight : ARM プロセッサに搭載されたデバッグ・トレース機能群。プログラムの実行経路を記録可能
- ・プロセッサトレース : CPU レベルでプログラムの実行経路を記録する機能。Intel PT、ARM CoreSight 等がある
- ・RISC-V : オープンソースの命令セットアーキテクチャ。IoT 機器等での採用が進む
- ・FPGA : Field-Programmable Gate Array。ハードウェア回路を再構成可能な集積回路
- ・PLC (Programmable Logic Controller) : 産業用制御機器。工場や重要インフラで使用される

<強化学習関連>

- ・強化学習：試行錯誤を通じて報酬を最大化する行動方針を学習する機械学習手法
- ・バンディットアルゴリズム：強化学習の単純化モデル。複数の選択肢から最適なものを効率的に探索する
- ・ミクロ最適化：単一ファジングアルゴリズム内の変異操作の適用順序・頻度を最適化すること
- ・マクロ最適化：複数のファジングアルゴリズムの選択・切り替えを最適化すること

<学会・カンファレンス>

- ・ACSAC (Annual Computer Security Applications Conference)：セキュリティ分野の国際会議 (Core Conference Rank：A、採択率約24%)
- ・NDSS (Network and Distributed System Security Symposium)：セキュリティ分野のトップ国際会議
- ・ISSTA (International Symposium on Software Testing and Analysis)：ソフトウェアテスト分野の国際会議
- ・HITCON：台湾で開催される国際セキュリティカンファレンス
- ・SECCON：日本最大級のセキュリティコンテスト・カンファレンス

<AI・競争環境関連>

- ・XBOW：2024年設立のAI駆動型自律ペネトレーションテスター。HackerOneで1位を達成
- ・Big Sleep：Google Project ZeroとDeepMindが開発したLLMベースの脆弱性発見エージェント
- ・LLM (Large Language Model)：大規模言語モデル。GPT、Claude等

<組織・制度>

- ・NISC (内閣サイバーセキュリティセンター)：日本政府のサイバーセキュリティ政策の司令塔
- ・重要インフラ：NISCが定める15分野 (情報通信、金融、航空、空港、鉄道、電力、ガス、政府・行政サービス、医療、水道、物流、化学、クレジット、石油、港湾)
- ・EDSA 認証 (Embedded Device Security Assurance)：制御システム機器向けのセキュリティ認証制度