



内容

- I. 第35回通常総会開催
- II. 第35回通常総会・特別講演
- III. 令和7年度役員紹介
- IV. お知らせ
- V. 編集後記

I. 第35回通常総会開催

第35回通常総会が令和7年6月18日(水)、品川区の貸会議室にて工藤副会長の司会のもとで開催され、盛況のうちに終了した。当日は52名出席(会場からの出席28名、オンラインによる出席12名)他、書面による出席12名のもと、会員数70名の3分の1(24名)以上の出席があり総会が成立していることが報告された。



工藤 副会長

1. 通常総会議事

(1) 阿部会長挨拶

「品質保証研究会の会長を務めさせていただいております、東京大学阿部でございます。本日は暑い中お集まりいただきありがとうございます。品質保証研究会はこれまでの活動を通じて実績を重ねてきました。講演会、見学会を重ね様々な知見を得る機会を得ました。また品質保証研究会の第1、第2グループとも成果も見えてきたという状況であり、今日の総会で内容について紹介をさせていただく予定です。会員の皆様におかれましては、忌憚なきご意見を頂ければと思います。」



阿部 会長

(2) 議長選任 会則/細則に従い、阿部会長が議長に選任され、(3)項議事が行われた。

(3) 議案審議

第1号議案: 令和6年度活動報告ならびに収支決算承認の件

錦野幹事、西田幹事より、それぞれ活動報告ならびに収支決算報告が行われた。また、藤巻監事より会計監査報告が行われ、両案ともに提案どおり承認された。



錦野 幹事



西田 幹事

第2号議案: 令和7年度活動計画ならびに収支予算案承認の件

錦野幹事、西田幹事より、それぞれ活動計画ならびに収支予算案の説明が行われ、両案ともに承認された。



藤巻 監事



沖田 幹事

第3号議案: 令和7・8・9年度顧問選任の件

沖田幹事より、顧問選任紹介が行われ、承認された。

2. 定例研究会活動報告

第1、第2グループ各々の令和7年度の活動等について各グループリーダーより報告された。

(1) 第1グループ: 鈴木リーダー

テーマ: 国内外最新知見を踏まえた品質コンプライアンス違反を発生しない/させないための QMS 等の研究

2017 年の大手鋼材メーカー製品に関する品質コンプライアンス違反が公表されて以降、製造業者による品質コンプライアンス違反が断続的に報告されており、これら品質コンプライアンス違反は原子力業界全体の信頼性にもかかわる重要な課題であることから、品質コンプライアンス違反を防止するための施策が必要となった。



鈴木リーダー

製造業者が陥りやすい品質コンプライアンス違反の早期発見、未然防止ができるように、かつ、各社の QMS に反映しやすいように ISO9001 要求事項と紐づいたものとすることを目標とし、原子力業界に携わる企業に対してユーザーフレンドリーな「国内外最新知見を踏まえた品質コンプライアンス違反を発生しない/させないためのガイド」を発行することとした。

品質コンプライアンス違反の事例を調査し、その抑止策を研究・集約しガイドとしてまとめるもので、品質コンプライアンス対応について「低減」「推奨」「助言」の提案レベルに分類し 26 件の事例について対応内容をまとめた。このガイドを会員企業の皆さんへ展開し、各社の品質保証活動の中に活用されることを期待するものである。

2025 年 8～9 月に「国内外最新知見を踏まえた品質コンプライアンス違反を発生しない/させないためのガイド」を品質保証研究会ホームページに公開する。

【質疑応答】

Q: 業種によってガイドの内容を噛み砕いて、各社の業態に合わせた対応を望むと説明があった。本ガイドはそもそもどういう業種での適用を想定しているのか。

A: 製造業種を意識している。工事(施工)業務などについては、その内容に応じてガイドを解釈して欲しいと思っている。

Q: コンプライアンス違反、すなわち倫理に反することをしてしまう背景は“うっかり”から、なにか“重い背景”があって違反をしてしまうまで様々な背景があると思う。背景が異なれば対策・対応も異なると思うが本ガイドを作成する際に背景は考慮しているのか。

A: どういう背景なのかというのは勘案してガイド化している。

Q: ルールを守らないという違反をする背景として“ルールを知らない”ということが原因であることもあり、知らないのであれば教育するというアプローチか。

A: NEA(原子力機関)で言われている特に日本人の特性では、“ルールを知らなくても仕事が進む。なぜなら【前と同じやりかた】を踏襲するから”というものがある。例えば計測器が進歩していて、進歩に応じてルール(方法)も変えなくてはならないはずだが、現実的でない従来のルールに固執していることがあり得る。教育が重要であるが“守れ”という教育だけでなく、背景を踏まえた対応が必要である。

(2) 第2グループ: 酒井リーダー

テーマ: NHK の実践(調達先評価、監査方法の改善含む)の研究

〔NHK: 無くす、減らす、変える〕

原子力サプライチェーンを維持・強化するためには、原子力事業に携わるサプライヤーに原子力に対してポジティブなイメージを持って積極的に活動していただくことが重要である。

サプライヤーに積極的に活動してもらうためには、サプライヤー自身に規制や要求事項を正しく理解していただくほか、発注者としてもポジティブなイメージの発信や業務量の低減を図っていく必要がある。これらを実現するため発注者の調達管理活動に着目し、“NHK: 無くす、減らす、変える”の取組方法を整理したガイドを取りまとめた。

ガイドは、調達管理活動を 11 のプロセス(*1)に分類し、NHK 事例について課題/問題点、改善内容/効果、NHK 実践に至るまでのアプローチや着眼点、注意点、満足すべき要求事項を示す解説を一件一葉にまとめている。



酒井リーダー

*1: 「①調達計画作成」「②調達先評価・認定」「③調達先監査」「④品質保証仕様書/購入仕様書作成発行」「⑤調達先の引合仕様書評価」「⑥調達先と製造開始前確認」「⑦製造中確認」「⑧立会検査(顧客立会含む)」「⑨製品リリース(出荷許可)」「⑩成果物の適合性確認(検収)」「⑪調達先不適合対応」

このガイドは品質保証研究会の会員が様々な業務プロセスの中で、NHK を念頭において継続的に改善に取り組む際の参考として活用されることを目的としている。各社がこのガイドを通じて要求事項を正しく理解し業務量を低減することが期待でき、ひいては原子力サプライチェーンの維持・強化につながると考えている。

2025 年 8 月に「NHK の取組方法を整理したガイド」を品質保証研究会ホームページに公開する。

【質疑応答】

Q: ガイドとして取りまとめる評価フローにおいて NHK 事例で未解決とされているものがあり、これが次年度の研究対象という報告であった。具体的には何か(25 年度に深堀解析が必要と記載されている「品質記録の保管」「計測器の管理」の 2 つが未解決ということか)

A: その通り。25 年度に深堀していく。

Q: 今回の NHK を進めることで、逆に管理を深める(増やす)ことが必要となるようなものはないか。

A: 単に「増える(増やす)」というより、変える(変わる)という形で管理の形が変化したものはある。

Q: NHK の中で一番多かった事例は何か。ガイドでは無くすという分類が無い(少ない)ように見える。

A: 多いのは「減らす」と「変える」という結果であった。「無くす」は「減らす」「変える」に付随する形で整理されたものと考えている。「無くす」だけに当てはまる事例は無かった。

Ⅱ. 第35回通常総会・特別講演

『AIが前提となる時代の品質・安全性へのアプローチ』

講師: 大学共同利用機関法人 情報・システム研究機構 国立情報学研究所
アーキテクチャ科学研究系 准教授 石川 冬樹 氏

1. はじめに

AIの品質、安全性といった観点での一般化した内容を説明する。

2. 機械学習型AIの品質・トラスト

この分野における機械学習とは、簡単に言えば「過去のデータから計算式を作る」というものである。評価対象の事象に、パラメータ(判断基準)を与えると「予測」ができる。予測＝計算式という図式である。

例えば $y=ax+b$ という一次式で計算式を作るのならば、 a と b という2つのパラメータを過去のデータから決めればよい。実際には何百万というパラメータで予測関数を表現することになる。「過去のデータをパラメータとして計算式を作る」というのが機械学習型AIと思ってよい。分類したいというタスクであれば、なにがしかの訓練データから分類の「線引き」を設定する・・・というのを行うのが機械学習である。とはいえその線引きがなぜ行われたか(人間は)分からない。

例えばテナガザルとパンダを「線引き」したことについて、なぜその境界で線を引いたかについて、例として目の周りが黒いとか言うことはできるけれど、それがすべてかと問われるとすべてでは無いようだし、そもそも目とはなんですか、周りとはどこですか・・・結局言葉にできない。ただ、例がいっぱいあると判断できる(線引きができる)ということだとすればそれはそれで便利だと・・・それがAIである。

さてこの分類に新たなパンダの画像を持ってきたら、ちゃんとパンダと線引き分類された側に配置されますか、ということが誰にも保証は出来ない。AIプロダクトはそういうもの。そういうものを我々は使おうとしている。

人が規則性を書きだすことが困難な機能であっても、機械学習させ識別・判断をさせることが可能である。機械学習のアウトプットは「演繹的推論」や「ルール処理」ではない。例えば翻訳であれば、主語がどれ、述語がどれというようなロジックで結論を導いているわけではなく、よって“判別や予測の性能の不確かさ”、“振る舞いの不確かさ”をはらんでおり、基本的にはデータに対する相対的な評価しかできないものであり、大量かつ「適切な」訓練データが必要である。

機械学習型AIシステムの品質は「AI品質」としてAIQM:機械学習品質マネジメントガイドライン(*2)、QA4AI:AIプロダクト品質保証ガイドライン(*3)で定義、運用している。(欧州では倫理という観点からのガイドラインがある)2020年以降は、倫理・トラスト・アライメントと言った用語で人間・社会観



<講師ご紹介>

国立情報学研究所
アーキテクチャ科学研究系 准教授
トップエスイープログラム 代表

総合研究大学院大学
複合科学研究科情報学専攻准教授

電気通信大学
大学院情報理工学研究科 連携准教授

ソフトウェア工学・特にディペンダビリティ:
形式手法・自動テスト生成・安全性論証など、および SE for AI & AI for SE、機械学習型 AI のエンジニアリング支援、産業界向け教育・実践研究 等
「スマートなシステム・スマートなディペンダビリティ保証」をビジョンに掲げ、組織を横断する研究室において活動されている

<http://research.nii.ac.jp/~f-ishikawa/>

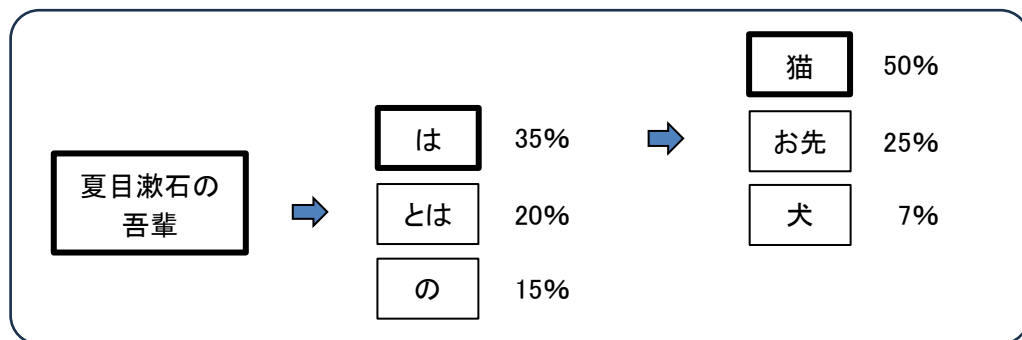
点に焦点があてられ、ISO標準や政府レベルのガイドラインなどが発行フェーズになっている。また、ChatGPTなどの対話型生成AIに対する議論に焦点が急速に移行している。

* 2: <https://www.digiarc.aist.go.jp/publication/aigm/>

* 3: <https://www.qa4ai.jp/download/>

3. 大規模言語モデル・生成AIによる変化

大規模言語モデル(LLM: Large Language Model)は自然な言語列を学習した機械学習モデルで、画像などの入出力を扱うマルチモーダル版にも発展している。LLMの仕組みは(雑に言えば)文節の次に来る文節を統計的な観点から「予測する」ことを繰り返していくもの。実際には「返答が人間にとって望ましいかという判断」を別途学習させるなどの工夫を加えて実用的なものとなっている。LLMを活用し自然な対話やタスク処理を行うChatGPT、Gemini、Claudeなどの対話型・チャット型生成AIが身近になっている。



ルールや知識に基づいて処理しているわけではなく、機械学習として、データから続きを予想しているものであり、例えば「トム・クルーズの母親は誰か？」に「メリー・リー・サウス」と答えられるが、「メリー・リー・サウスの息子は誰か？」に答えられない。メリーの息子を問う人(入力する人)がいないので、学習しない。三段論法のようなロジックで答えを導いているものでなく、単語の次の単語を予測(学習結果)している。

また、LLMは訓練データのカバー範囲や不適切なバイアス(例: 看護師と言えば女性)に影響を受けるという課題がある。基盤モデルとして、誰かが汎用性の高いモデルを作ってくれていれば、誰でも使えるという利点がある。画像の見方というのは学習している(基盤モデルとしてある)ので、「我社に特有の不良品の見方を教え込めば(楽に)使える」というような使い方ができる。

「一人一人がうまく使おう」とすればどんどん進む(使える)。

チャット型生成AIの品質・受入という観点で世界が大きく変わった。大半の組織の活動において、データの収集や訓練をせず、カスタマイズと評価が活動の中心になっている。入出力が自由で、想定すべき入力が多様多様になり、出力の判定も自動化がしづらい・曖昧なものになってしまう。試す・動かすが簡単に進む一方で「あなたの仕事ならではの品質が実現できたか？」という自分たち固有の評価や保証には大きなコストがかかる実態になってきている。

AIにおけるソフトウェア品質として、QA4AI:AIプロダクトの品質保証(Quality Assurance for Artificial-Intelligence-based products and services)において(下表の)品質特性を明確にしている。

回答性能	・ 求められるタスクに対応する 副特性: 自然言語処理、ツール活用、創造性・多様性、制御可能性
事実性・誠実性	・ 検証可能な事実・知識に即した回答をする(不明なことを断言しないことを含む) 副特性: 一般的な事実に対する・・・、与えた事実に対する・・・、根拠の説明性・妥当性
倫理性アライメント	・ 人間や社会に対して害をなさない適切な回答をする 副特性: 公平性、安全性、データガバナンス
頑健性	・ ノイズや想定外の入力に対して品質を維持できる
AIセキュリティ	・ 悪意ある入力に対して品質を維持できる
その他考慮事項	・ 透明性、説明可能性、アクセシビリティ、ユーザビリティと社会心理的側面、機能適応性

これらを踏まえAIの評価は大変。問いかけをどうするか。犬か猫か、○か×かはまあ判定可能であるとして、人によって評価が分かれるなどの判定はどのようにすべきか・・・。

「我社の業務として・・・」という答えを求めるのであれば、それが正しいかのオリジナル:個別ケースに応じたテストが必要である。

4. 品質保証の道具としてのAI

LLMのおかげで「なんでもできる」こととなったが、一から十までLLMでやらせることがベストとは限らない。雑な書き方(評価)をLLMにやらせて、詳細を別でやるという手段もある。

例えばソフトウェアを評価するとしても、何でもAIにやらせるというわけではなくて、既存の技術予報が強いということもありえるし、過去のデータばかりでなく、試して確認(トライアンドエラー)で確認していくことが有効である場合もある。LLMでなんでもやらせればいいというわけではない。結局テストをするAIをテストするAIが要・・・のようなことになる。AIを道具として使うということでどう評価するかを考えていく必要がある。

アルゴリズム・論理推論として、タスク特徴のアルゴリズム、モデルベース、制約ソルバーなどの論理推論技術をテストの支援として用いる手法がある。ここで、倫理的なルールを守る・定量指標を上げるなら従来技術(とLLMを組み合わせたもの)が強い可能性はある。

データから機能を獲得する「教師あり」学習技術の活用(画像から部品の自動抽出など)や、進化計算や強化学習などの探索・最適化技術の活用をテストの支援として用いる手法もあり、LLM以前のAIに対してはこれを使用していた。特定の定量指標を上げるなら、(やはり)従来技術(とLLMを組み合わせたもの)が強い可能性が高い。

ところで、AI界限は変化が早く、追いきれず、「確立したことを学び続ける」だけではAIの活用には不十分である。適切にAIを使うには、AIへの指示・AIの出力の評価における原則・技術を持つことが不可欠である。

5. まとめ

- ・AIの活用は試行錯誤が大事な世界である。
- ・不完全さ・不確実性が原則にある。
- ・ステークホルダーとの議論を通じて、不完全・不確かなものを受け入れ、自動化された探索・テスト・測定・監視による継続的進化を目指している。
- ・新しい技術・時代を楽しんでいきましょう。

【質疑応答】

Q: AIを活用しようと考えており、気になるところはセキュリティ(守秘)であると考えています。AIの環境(データベース:DB)等が重要であるが、DBは多くのデータサンプリングのために外に開かれるネットワークにつながるように構築することが望ましいと考えるものの、強度に守秘をする場合には、独自にDBを持つことになると考えています。先行例(保険会社等)は、どのように守秘を管理しているか承知していれば教示願いたい。

A: 保険の規定は公開されている。社内の規定やマニュアルはデータをクラウドに置くことには無理な場合があることも事実。LLMとDBを自前サーバに置き守秘を確実にしている。サーバの計算性能が限られ、ChatGPTのような大規模なモデルを動かすことはできないため、性能が良くない(下がる)ことが考えられる。

Q: AIの活用について、コンサルティングはあるのでしょうか。

A: 内容による。AIの正しさ、評価の妥当さは、自分自身でも普段意識していない領域の「こういうものとこれを照らし合わせている」や「これが評価基準だ」というのは指示ができる。知識としては、それを言語化できるか。指示が明確であればできる。..それが大事。

Q: ChatGPTなどのAIを使っていくというのは、正しさをどのように検証していくかが重要だと思っています。人間の正しさに留まらず人間を超えるのか。AIの将来性についてご教示願いたい。

A: 得意なところが違う。同じ時間でできるコーディング(仕事)の量は圧倒的にAI優位。新しいものを作るというときに、過去の知識を取り入れることであればAIは得意である。それ以外の部分では「AIができていて」ことを人間が評価できないのでAIが優れているかわからない。とはいえAIはツール以上のものではないと私は思っている。

Q: AIが正しくない答えを出してしまうこともあり得ると思います。どうすれば改善できますか。

A: 公平性を上げると正解率が下がるなど、過去の人間の選択が間違っていることもある。公平性という評価も結果の公平性と手続きの公平性がある。個人の意見は社会に対してどうこうというものではなく、あるべき姿というなかで複数のゴールがある。結局のところ正しさは経営判断になるかもしれない。

Q: AIが「正しくないことを出力する」ことを知らないと信じてしまうかもしれないと思っており、どうすれば気づけるかという点でアドバイスを頂きたい。

A: AIを使う上で、「いいプログラムとは何か」ということを教え、検証しなければならないということに尽きると考えている。

Q: AIを使うことはリスクがあると思っています。「リスクがあるが、リスクを受け入れろ」は欧米的な発想で、アジアはリスクを受け入れない＝安全は無料(コストを払わない)がまかり通るなかで、リスクの管理についてご意見を頂きたい。

A: 「リスクゼロはないよね・・・」は、実際、「こんな間違いするAIは使えないよね」という話はある。日本の文化としてそういうところがあるよねと。とはいえ日本人は、「一度受け入れると気にならなくなる」という特性があると思っている。このことから、リスクに対して日本人は敏感だということはあるか知らないが、日本とヨーロッパが結局のところ、リスクゼロはあり得ない、80%を受け入れなければならない、リスクはあるが進めざるを得ないということが現状。

Ⅲ. 令和7年度役員紹介

1. 役員紹介

会 長	阿部 弘亨	東京大学大学院工学系研究科原子力専攻・教授
副会長	工藤 竜太	東芝エネルギーシステムズ株式会社
副会長	芝原 啓介	株式会社 日立製作所
副会長	宇奈手 一之	三菱重工業株式会社
幹 事	日隈 聡	東芝エネルギーシステムズ株式会社
幹 事	高橋 淳	株式会社 日立プラントコンストラクション
幹 事	西田 徹	GE日立・ニュークリアエナジー・インターナショナル・エルエルシ
幹 事	高次 正弥	三菱重工業株式会社
幹 事	錦野 嘉浩	株式会社 日立製作所
幹 事	沖田 康典	三菱重工業株式会社
監 事	藤巻 真吾	株式会社グローバル・ニュークリア・フュエル・ジャパン
監 事	田上 貴吉	一般社団法人 原子力安全推進協会

2. 退任役員

精力的に第1グループ活動に取り組み、また幹事としても長年ご活躍された清宮 明氏(東芝エネルギーシステムズ 株式会社)におかれましては、突然のご逝去により退任することになりました。これまでのご尽力に深く感謝すると共にご冥福をお祈り申し上げます。

IV. お知らせ

第35回通常総会において報告しました研究会成果である以下①、②のガイドを、品質保証研究会のホームページにて公開致しました。会員の皆様におかれましては、各社の品質保証活動の中で活用していただければ幸いです。なお、これらガイドの公開は会員の皆様に限定させていただきますことをご了承ください。

- ① 「国内外最新知見を踏まえた品質コンプライアンス違反を発生しない/させないためのガイド」
- ② 「NHK の取組方法を整理したガイド」

V. 編集後記

生成 AI の規制議論が世界的に進む中、日本でも産業界での活用とリスク管理のバランスが問われています。そんな中での今回の特別講演は、AI を「使いこなす力」の重要性を改めて感じさせるものでした。また、電力業界では GX(グリーントランスフォーメーション)や再エネ導入の加速が進む一方、原子力の役割も再評価されつつあります。こうした変化の中で、品質保証の視点がますます重要になっていることを、本会の活動報告からも実感しました。変化の時代だからこそ、地に足のついた品質活動を心掛けたいと思います。

(編集:A.H)

以上