

EASTICA 発表 – 2024 年 11 月 12 日 – John Sheridan

はじめに

スライド 1

こんにちは、英国国立公文書館デジタル部門ディレクターのジョン・シェリダンです。今日は皆様とご一緒できて大変光栄です。

スライド 2

これから、人工知能（以下、AI）とデジタルアーカイブズについてお話します。

ここからの約 45 分間、皆さんが理解を深められるように、このプレゼンテーションは 3 つのパートに分かれています。

最初のパートでは、何が AI で何が AI でないのかをお話します。

パート 2 では、AI をどのようにアーカイブズに利用できるかを考える。AI が解決に役立つ問題にはどのようなものがあるだろうか。

パート 3 では、AI の導入に関するより広範な問題、つまり実用面、法律面、倫理面、環境面の考慮事項について説明する。

皆さんと一緒に、現代に取り組むべきさまざまな問題を考えていきたい。

スライド 3

その前に、まず私についてお話ししよう。私は英国国立公文書館（以下、TNA）のデジタル部門ディレクターに就任して 9 年になる。TNA には 2006 年から勤務している。公務員としてのキャリアがあるが、アーキビストになったのは全くの偶然である。

すぐにわかるかとは思いますが、私の専攻は数学とコンピューティングである。

子どもの頃から、コンピュータやその機能に興味があった。初期の家庭用パソコンから現在の AI になるまでの進化を目撃できたことは、非常に興味深い経験であった。

振り返ってみても、私は特に先見の明があったというわけではない。

今、私たちが生きている世界は、私が期待していた、あるいは想像していた未来ではない。

もう一つわかっていることは、自分たちは、より良い方向のために影響を与えることができるということである。アーカイブズやアーキビストには、この新しい AI の時代に、重要な役割があると信じている。

実際、私たちは、この革命に積極的に関わる責任があると思う。アーキビストの視点、知識、スキル、感性は、フロンティア AI (Frontier AI) において相当に必要とされている。

アーカイブズで働く当事者たちは、まだそのことに気づいていないかもしれない。つまり、これまで以上に、自分たちがやり慣れている仕事に関心を持つ必要がある。

私たちは、コレクションを積極的に武器にしていくべきである。また、アーカイブズ専門職の国際的なコミュニティとして、この取組に力を合わせていくべきである。

第1章 – 人工知能

スライド5

私は生まれながらの楽天主で、コンピュータが大好きだ。当然ながら、私はAIの登場にとってもワクワクしている。コンピュータを使って、今何ができるのか、そして、今後何ができるようになるのか、その可能性を考えると、なお一層ワクワクする。

しかし、楽観視できないこともある。

もちろん、純粋に優れたものであり、最小限のマイナス部分しかないというイノベーションもある。ポリオや麻疹の予防接種など、多くの医療にかかわる措置がこのカテゴリーに属する。しかし、メリットとデメリットのバランスがどこにあるかを把握するのが難しいイノベーションもある。

30年前、私はワールドワイドウェブ（www）の出現にワクワクした。私は、それが人類にとってよいものをもたらすものだと思っていた。そして、その考えは長年変わらなかった。今は、その確信が持てない。ウェブが私たちにもたらした利益は計り知れないが、それと同じくらいの弊害も簡単に挙げられる。

大きな教訓がここにある。技術の利益と弊害が、技術自体に内在することは、まれにしかない。問題は、我々が技術を使って何をするかということにある。

私たちが主体なのである。

スライド6

イノベーションの普及理論によると、十分に有用な技術はいかなるものでも時間の経過とともに、広く採り入れられていくのは避けられない。コストが低下することで、新しい技術は浸透していく。これは、特にさまざまな用途に活用できる、汎用技術においてよく起こることである。

では、新しい技術を発明する人々はどうであろうか？ 彼らは、自分たちが開発したイノベーションの受益者になることはほとんどなく、イノベーションの用途を最終的に決定することもない。

スライド7

ヨハネス・グーテンベルクは、聖書を効率良く複製したいと思っていた。彼が発明した活版印刷術は、数学や科学を含め、さまざまな革命のきっかけとなったが、それは彼が想像さえしていなかったし、望んでいたことでもなかった。

スライド 8

では、AI に関しては、どの程度楽観視してよいのであろうか？

先月公表された「愛のマシン：AI はより良い世界を作れるか (Machines of loving grace, how AI could transform the world for the better)」というタイトルのエッセイの中で、AI 分野をリードする企業である Anthropic 社の CEO、ダリオ・アモデイ氏は、「AI がいかに劇的な影響をもたらすか、ほとんどの人が過小評価していると思う」と述べている。

彼はさらに、すべてがうまくいった場合、強力な AI がどのようなものになるのか、そして人類に何をもたらすのか、詳しく説明している。

アモデイ氏の挙げる進化の例は、生物学と身体健康、神経科学とメンタルヘルス、経済発展と貧困削減、平和と統治、そして最後に仕事と意義にまで及ぶ。

スライド 9

〔人工知能プログラムの〕AlphaFold を開発者した、デミス・ハサビス氏とジョン・ジャンパー氏は、AI を活用してタンパク質の構造を予測する研究でノーベル賞を受賞した。この発見は、最終的に癌の治療法を見つけるための第一歩となるのかもしれない。

懐疑的な人でさえ、コンピュータが大きく進化し、つい最近まで処理できなかったあらゆることができるようになったことを認めている。

スライド 10

最も注目すべきなのは、いわゆる生成 AI が新たな時代に入ったということである。生成 AI とは、プロンプトに応じてテキスト、画像、動画、コンピュータ・コードを生成できる AI である。

これは、ディープラーニングとも呼ばれる、ディープニューラルネットワークの開発が劇的に進化した成果である。

私たちの多くが、スマートフォンには ChatGPT、職場では Co-pilot を導入しているが、どちらも素晴らしいツールである。私たちは、自然言語プロンプトを使用してシステムにあらゆる質問を投げかける。必ずしも正しい答えではないが、一貫した答えは返ってくる。

ただし、一貫していることと正しいことは違う。そこに問題が生じるのである。

それは違い以上のものであることを、これから詳しく説明する。

スライド 11

私たちは、どうやってこの技術にたどり着いたのであろうか。そして、この技術がこうした特徴を備えているのはなぜだろうか。

現在の AI の基盤となっている数学的な概念は、古くから存在している。科学や技術と関連する数理においては、よくあることだ。アインシュタインが「特殊相対性理論」を発見する半世紀前にあたる 19 世紀中頃、リーマンはリーマン曲率テンソルを考案した。アインシュ

タインは〔特殊相対性理論に〕リーマンの数理をかなり必要としていたが。

AI の数学に関して言えば、線形代数（ベクトルと行列）も 19 世紀の間に開発がかなり進み、1920 年代後半までにはほぼ完全に定式化され、現代の私たちが理解しているものに到達していたことがわかっている。これは、万能チューリングマシンや、私たちが現在理解している計算可能性理論の基盤となったチャーチ＝チューリングのテーゼよりも以前のことであり、そして現代の「情報理論」と呼ばれる、シャノンの「通信の〔数学的〕理論」よりも前のことである。

つまり、今の世代の AI の作成に利用される主な数学的手法は、極めて古くからあったものである。ほとんどの場合、コンピューティングやネットワーク通信の基盤となる概念よりも、はるか以前から存在していたのである。

スライド 12

人工ニューラルネットワークを作る概念も、かなり古くから存在する。

ウォーレン・マッカローとウォルター・ピッツは、1943 年に「閾値論理ユニット」の論文において、生体ニューロンのシミュレーションを提案した。彼らの人工ニューロンモデルは、重みを掛け合わせた複数の入力を受け取り、合計値を出して結果を生成する「関数」である。ニューラルネットワークの基本構成要素であるパーセプトロンは、1957 年 1 月、コーネル航空研究所のフランク・ローゼンブラットによって考案された。

スライド 13

多くの現在の AI の基盤となるのは、層状に配置された人工ニューラルネットワークである。これは、生体ニューロンやニューラルネットワークに着想を得たものであるが、形式を大幅にシンプルにしたものである。

このシンプル化により、行列を入力と重みに使用でき、線形代数の手法を使用して、行列を効率的に計算できるようになっている。

しかし、その過程で、生体ニューロンやシナプスの重要な特徴を見逃しているという可能性がある。つまり、それが現代の AI が抱える困難の原因なのかもしれないのである。

つまり、私が言いたいのは、現代の AI が強力なのは確かだが、魔法ではないということである。単なる技術である。そして、ほとんどの部分は、100 年前の数理を使って理解することができるのである。

スライド 14

ここである疑問が湧いてくる。数学や手段でなければ、現在話題を呼んでいる生成 AI をもたらした飛躍的な進化とは何だったのであろうか？

簡単に言えば、つい最近まで、私たち人間には膨大な量のデータ演算を行う能力がなく、また、膨大な量の演算を行うためのデータもなかった。

膨大な演算能力と膨大なデータは、現代の生成 AI の前提条件である。

現在の大規模言語モデル (Large Language Model。以下、LLM) は、ワールドワイドウェブ全体のデータを使用してトレーニングされている。何兆語ものテキストを取り込み、演算能力とそれに費やすお金の両面で、多大なコストをかけてきたのである。

スライド 15

LLM のヒューマンインターフェースとなるのがチャットボットである。ここではじめて、まるで人間同士で会話しているかのように、自然言語を使用することでコンピュータとコミュニケーションをとることができるようになった。

旧世代のチャットボットとは異なり、ChatGPT などの LLM のチャットボットは、知能があるように見える。

少なくともスタンフォード大学人文科学部の研究者によると、最新バージョンの ChatGPT は、厳格なチューリングテストに合格したということである。

ここで特に重要なのが、自分の批判的な視点を維持することだ。

AI をうまく活用するには、現在の AI に何ができて何ができないのか、そしてその理由を、私たちは理解しようとしなければならない。

これは、簡単なことではない。

スライド 16

生成 AI システムの開発に携わる企業にとって、そのリスクは非常に大きい。LLM の制作には莫大なコストがかかり、ユーザーの要求に応えるためにモデルを実行する場合も、同じようにコストがかかる。すべての資金を費やしても、その上、コストは増え続けているため、投資家に対して利益を還元しなければならないというプレッシャーを感じるのは当然理解できる。したがって、積極的に製品を売り込みながら、おそらくその限界については曖昧にしておくのである。

大手テクノロジー企業は、語り手として信用できないとだけ言えば十分であろう。

スライド 17

すべてを理解するのが難しいその根幹は、知能の定義が曖昧だからである。

もちろん、私たちは知能が何かを知っているつもりである。私たちは、人間の脳という最も素晴らしい生物学的知性で得られる経験を通して、直感的に理解をしている。

しかし、知能に関しては、広く受け入れられている科学的理論は存在しない。実際、心理学者たちは、数十年にわたって、このテーマに関する議論を続けている。

ある状況でどの種の知能が発揮されるのか、あるいは、実際、あるタスクを実行するのにどの程度の知能が必要なのかを把握しようとするにあたっては、私たちの直感はあまり役に立たない。

私たちがほとんど何も考えずにできる、私たちにとっては簡単な動作、例えば、コップ1杯の水を1口飲むという動作だけでも、驚くほど多彩な知能が求められるものだ。

スライド 18

生物の知能の形は、数億年にわたって進化を遂げてきた。

人間の脳のニューロン間の接続は、ウェブ全体でトレーニングされた人工ニューラルネットワークのニューロン間の接続よりもはるかに複雑である。

一般的な人間の脳には800億個以上のニューロンがあり、100兆個以上のシナプス、つまり接続がある。一方、自然界では、化学的、電気的の両方のメカニズムを用いて、多種多様な組み合わせで情報を伝達している。

ここで重要なのは、人間の脳を過小評価したり、人間の知能が特定のシステムを作ることに貢献しているのを軽視したりしてはいけない、ということだ。

スライド 19

人間の脳は、他の高度に洗練された知能、つまり、他の人間が多数存在し、複雑で混沌とした予測不能な世界を生き抜くために進化してきた。私たちは、人間の世界を理解するように作られているので、物事を擬人化する。私たちは、物質や生物がまるで人間であるかのように、自然に、そして簡単に、動機、意図、感情、さらには意識さえも彼らに持たせる。

AIとは何か、それが何に役立つかを理解しようとする、私たちの直感は私たちを誤った方向に導くのである。

商業上の動機で利益を誇張したり、実際の状況を偽装したり、あるいは背景を隠ぺいしたりすると、その困難さは深刻化する。

スライド 20

AIの限界の具体例をいくつか挙げてみよう。AIをうまく導入するには、これを理解することが重要である。

こうした限界に対する認識不足が、AIプロジェクトの失敗によって、今後1、2年のうちに多くの失望が生まれる根本的な原因になるのではないかと思う。

スライド 21

まず、ハルシネーション（幻覚）である。

ChatGPTなどのLLMは、一見それらしく見えるものの、実際は誤ったテキストを生成する。

ハルシネーションが起こるのは、LLMが、確率に基づくテキスト抽出機であり、現実世界のモデルとプロンプトに対する応答を評価するのに必要な批判的思考能力が欠けているためである。

LLMは確率に基づいているため、ハルシネーションは一見それらしく聞こえる。それが危

陰なのだ。LLM が出す回答は、それが完全に事実と異なっていたとしても、おそらくは、良い回答である。

ハルシネーションは、確率の高い引用やテキストを生成することで軽減されるが、LLM の仕組みを考えると、LLM が事実を把握する上で直面する根本的な問題を現在の技術を使用してどのように克服できるのかは、予想しがたいものである。ナレッジグラフは役に立つのであろうか？

重要なのは、事実を知るためには、LLM に頼らないことである。

これは LLM が知能の一種として、あるいは事実の「候補」となる情報の最初のソースとして有用であることを否定するものではない。しかし、LLM をこのような方法で使用し、事実を知ることが重視する場合は、その後のワークフローに「事実確認」タスクを追加する必要があるだろう。

スライド 22

2 つ目は、破滅的忘却である。

これは、人工ニューラルネットワークに基づくすべての AI に共通する問題である。一度トレーニングすると、ニューラルネットワークに新しい情報を入力し、過去に学習したすべての情報を維持しているかどうか、確認のしようがない。新しい情報でモデルを更新すると、それまでに学習した内容を忘れてしまい、しかも予測できない形で問題が発生する。

つまり、コンピュータとしてコストが高く非実用的な ChatGPT などのデータセット全体でモデルを再トレーニングすることなしに、モデルを段階的に改善するのは、実務的には不可能であるということだ。これが、OpenAI が GPT の基礎モデルを毎日少しずつ改善するのではなく、一定の間隔でリリースしている理由である。

破滅的忘却を管理する手法は確かに存在する。例えば、ChatGPT はモデル外部の「保存された」情報に事例としてアクセスできるが、人工ニューロンの現在のシンプル化されたモデルに固執する限り、基本的な問題はおそらく残る。

スライド 23

3 つ目は、全く新しい状況や予期せぬ状況に対処できないことである。

現代の AI は、トレーニングデータ以外のものに直面すると、うまく処理することができない。しかし、私たちが生きているのは開かれた混沌とした世界である。これは、自動運転の車の開発者から見て、特に重要な問題である。AI の「ワンショット学習」能力の向上は、現在の研究における重要な分野である。

身近な例としては、デジタル記録の利用審査プロセスに AI を導入することを想像してほしい。AI は、たとえばまったくの新規のイベントなど、まったく前例のない機密性の高い情報の特定よりも、すでに問題がよく知られた機密情報の特定の方がはるかに優れているか

もしれない。

スライド 24

4つ目は、他の知能〔をもつ個体〕が何を考え、何をしようとしているかを予測できないことである。

私たち人間には、心理学者が「心の理論」と呼ぶものがあり、人間の脳は周囲の他者の精神状態を自然にシミュレートする。他者が何をするかを自分で予測する際の助けとするために、私たちの脳には他者が何を感じ、考えているかを示すモデルがある。私たちは、他の知能〔をもつ個体〕の精神状態をシミュレートする。これは驚異的な能力だが、私たちはごく簡単にこの作業をこなしている。

現在の AI には「心の理論」が備わっていない。このように理解力に根本的なギャップがあるため、LLMは、複雑な社会的相互作用において、文脈に合った適切な応答を生成することができない。

意思決定に AI を導入する場合、「心の理論」が欠如していることを踏まえておくことが重要になる。

現在の AI は、他者の意図や動機を推論することはない。AI には共感する能力はなく、対象となる人々が意思決定に対して抱くであろう感情的な反応や対応を認識しない。

だからこそ、意思決定においては、「心の理論」を備えた高度に洗練された知能、つまり人間が主導権を握ることが重要なのである。

パート 2ー アーカイブズにおける AI の利用

スライド 25

アーカイブズは情報機関である。これまで述べたように、AI は驚くべき機会をもたらしてくれる。

AI のアーカイブズへの応用の仕方は5つあることを、今日はお話したいと思う。AI に求められることの難易度がますます高まっており、そのため導入が成功する難易度も高まっている。

スライド 26

まず、データの拡充とリンクへの AI の活用である。

自然言語処理ツールは、古くから存在してきた。テキスト中の固有名詞を特定するために、従来はトークナイザーや規則、地名辞典検索を組み合わせ使用してきた。

エンティティとその関係の特定、記録のコンテキスト化は、Records in Context (RiC) などのイニシアチブの目的の1つである。AI は、既存の記録の記述を、この新しいフレームワークに変換するのに役立つだろう。

スライド 27

「私たちの遺産、私たちのものがたり (Our Heritage Our Stories)」研究プロジェクトでは、コミュニティが作成した遺産コンテンツをより伝統的なアーカイブコレクションと統合し、人々が探求できるようにする試みを行っている。私たちの研究では、AI を活用してコンテンツの充実を図り、人物、場所、イベントなどの主なエンティティを特定し、Wiki データのレファレンスを使用して、その概念を記述全体にわたってリンクさせてきた。

このプロジェクトは、この3年間継続して実施された。当初は「従来型」の自然言語アプローチを使用していた。当然、この取組は非常にスムーズである。

注目すべきは、LLMを使用するようAIパイプラインを作り変えることが比較的簡単だったことである。

直後から、LLMは、従来の固有表現認識アプローチと同等、またはそれ以上の結果を達成することができた。

スライド 28

2番目に、手書き認識、文字起こし、翻訳へのAIの活用である。

これも新しい話ではない。2016年から、手書きのテキストを処理し、書き起こしする「Transkribus」のようなツールは登場していた。

このような作業をこなすコンピュータの能力は素晴らしいものである。私たちは記録をデジタル化し、自動で文字を書き起こすことができる。こうした文字起こし機能は、検索エンジンで利用できるようにすることで、手書きの記録を発見しやすくしたり、また、まさにAIで活用したりすることもできるようになる。

また、AIを活用して、記録や記録の記述を別の言語に翻訳することもできる。コレクションについて国際的に取り組む際の言葉の壁は、急速に消えつつある。

アーカイブズのユーザーの中には、すでにこれらのツールを活用している人がおり、閲覧室で記録をデジタル化し、ほぼ同時に文字起こしと翻訳を行っている。この技術は一般的になりつつあり、今後ますます広がっていくであろう。ただし、このことは同時にアーカイブズに対していくつかの疑問を投げかけている。特に、ユーザーが新しい記録を自らデジタル化し文字起こしをする場合の、データ保護や著作権の問題である。

スライド 29

3番目に、コンテンツの要約にあたってのAIの活用である。

LLMは、単語を処理する能力が非常に優れている。特にボーンデジタルまたはデジタル化された所蔵資料について、LLMを試しに活用することで目録記述を新規に作成したり向上を図ったりするのは、アーキビストとしては当然のことだ。

また、目録記述自体を要約することも可能である。私たちは、AIを活用して、目録にある

記録の記述から短いテキスト・スニペットを作成し、ユーザーに目録記述の概要を簡潔に伝える実験を行っている。

この実験で得られた結果は、実装に耐えるものとして、十分に信頼性がある。

AI によって生成され、要約されるケースが増えるにしたがって、アーカイブズ目録はどのような位置づけになっていくか——これは興味深い問いかけである。

目録の重要性は今後も変わらないだろうが、その主な機能は変化するだろうというのが私の予感である。目録はユーザーの検索を補助するものではなく、より AI のリソースになっていくだろう。

目録が持つ記録の出所に関する情報、つまりコンテキストこそが、アーカイブズの真の強みである。

合成コンテンツ、そしてまさに偽造コンテンツがますます氾濫してくる中、ある記録が何の証拠となっているのか、誰が作成し、いつ、なぜ、どのようにその証拠が作成者によって当時行われていたことと関連するのかわかることは、ますます価値のあることになっていくであろう。

アーカイブズ目録は、ナレッジグラフとして再構築され、これから何年にもわたって、アーカイブズの中心であり続けるだろう。

スライド 30

4 番目に、調査をサポートし、可能とする AI アーカイブズアシスタントである。

アーカイブズの調査は難しいものである。アーキビストによるコレクションの整理・分類方法は、利用者にはわかりにくく、混乱する。目録の使い方を学び、アーカイブズに何があるかを探求するには時間がかかる。

時間や情報量の消化という点で、研究者の処理能力は限られている。数か月間かけてアーカイブズをすみずみまで調査する余裕のある人はわずかである（一部にはそれを実行する人もいて、その経験は非常に有意義だとされていることは理解している）。

AI は、研究者の処理能力を飛躍的に高めることができる。アーカイブズ調査アシスタントのチャットボットは、ユーザーが閲覧した情報をすべて追跡し続け、関連する検索結果やまだ詳細に閲覧していない記録の記述を追跡し続けることが可能である。また、調査の補助や学術論文の処理もするかもしれない。

自分で検索を実行するよりも、AI アシスタントに質問しやり取りして、調査を進めながら、AI アシスタントを構築することができる。

ここで述べたほとんどの技術、つまり ChatGPT を使用して作成したパーソナルな AI リサーチアシスタントは、すでに存在している。

こうした技術の評価は簡単なことではない。

LLM は、目録に対するクエリからの結果セットをどの程度要約できるであろうか。確かに、

結果の要約を意図して、テキストを生成することはできるが、研究者の動機、目標、クエリ、結果について、モデルは一体何を理解しているのであろうか。ユーザーの過去の研究分野や内容については言うまでもない。

これは、さらに大きな疑問を投げかける。研究分野におけるアーカイブズに対する人間の研究者の理解にとって、アーカイブズ研究の取組は、どれほど重要なのであろうか。検索結果や目録記述を精査することは、学習プロセスの一部である。すべての結果セットは研究者にフィードバックを提供する。私たちは、研究者のプロセスにおける重要な手順を自動化しようとしているのであろうか。

試してみなければ、それはわからないであろう。

スライド 31

5 番目は、デジタル記録の評価、選別、利用審査に AI を活用することである。

保存する記録を決定し、アーカイブズに移管するプロセスは、デジタル時代にはうまく対応できていない。

デジタル記録の評価と選別における AI の活用には、大きな可能性があると考えている。これについては、過去にガイダンスを公開したが、すでに AI の進歩に追い越されている。

私たちは、複数の政府機関のナレッジおよび情報管理チーム、および中央デジタルデータオフィスと連携し、保存〔期間〕、最終処分、評価の決定に AI を活用するための大規模な概念実証を構築している。

最終的な意思決定を行う人々にとって本当に役立つツールセットを見つけるため、さまざまなアプローチを試している。

簡単なものもある。共有ドライブにある、その他の種類のファイル（ログファイル、プログラムファイルなど）から「情報コンテンツ」をソートすること。記録全体からの固有表現抽出と結果のコンパイル。記録の対象となる人物、場所、組織、イベントの要約などである。より難しいがまだ取り組みやすいのは、「個人用」と「業務用」のメールを振り分ける作業がある。

同様に、永久保存の価値があるデジタル記録セットの選別後、これまで取るに足らないとされていた記録の中に類似するものがあるか、選別の対象とすべきかどうかを、AI を活用して確認するなどもある。

最も難しいのは意思決定である。記録が何の証拠となるのか、またその記録は保存する必要があるのか、の決定だ。当時、記録作成者は何を考え、何をしていたのか。将来、人々は何を知る必要があるのだろうか。

現在の AI の限界として知られるものが、ここにある。つまり、当事者の意図を理解する能力に欠けているということだ。

評価と選別の意思決定〔の範囲〕を拡大していくには、AI が不可欠になるであろう。しかし、レコードキーピングの意思決定者がこうした新しい機能をうまく活用するには、新しい

情報環境と新しい業務プロセスを作り上げる必要がある。

スライド 32

これまでは、AI を活用してアーカイブズの業務を支援できるかもしれないことを中心に語ってきた。

AI 導入においてアーカイブズが関心を持っている分野は他にもあり、これについても掘り下げる価値がある。

それは、一般的に意思決定を行うための AI である。

良くも悪くも、直接的、あるいは広範囲なプロセスの一環として、意思決定に AI が利用されるケースが増えている。レコードキーピングの観点から見て、私たちは何を保存すべきなのであるか。LLMを基盤とする、意思決定システムの記録をアーキビストは、どのように保管するのであるか。おそらく、利用可能なあらゆる手段を使って記録をキャプチャーするウェブアーキビストのアプローチは、巨大なシステムをキャプチャーし、管理する方法よりも優れたガイドとなるはずである。

AI システムの保存実務

AI ベースのシステムの何を保存するのであるか。証拠として保存するには、ニューラルネットワークについてのどのようなコンテキスト情報が必要であるか。AI の長期保存に関する新たなアーカイブズの手法の開発が急務である。

合成コンテンツ、情報の出所、信頼性

私たちの周囲には、出所不明の合成コンテンツがあふれている。経済や社会が必要とする、出来事の証拠はどこにあるのであるか。出所情報の技術には重要な進展があるものの、テキストではなく、画像や動画などのメディアコンテンツにより重点が置かれている。しかし、テキスト生成マシンは野放しの状態で、私たちの周囲にあふれかえっている。C2PAなどのイニシアチブが拡大し、テキスト資料に適用されることを期待する。

パート 3-AI の導入に関するより広範な問題

スライド 33

AI は、あらゆる場所で仕事のあり方を変えている。

アーカイブズとアーキビストは、私たちが保持し、管理する情報の、信頼できる管理者になることを目指している。AI の活用において、人々の信頼と信用を維持する必要があるのは、私たちだけではない。

つまり、責任を持ってAIを活用するとは、実際には何を意味するのであるか。

仕事にはどのようなタスクが含まれるのか、そのタスクのうちのどれが、全体的または部分的にAIの活用に適しているのか、慎重に検討する必要がある。

理由はすでに話してきたように、人間の知能を個別の知能の集合体に分解することはでき

ないが、業務をタスクに分解することはできるので、業務に AI を活かせる分野について理解を深めている段階である。

進化を遂げるため、私たちは業務、タスク、スキル、能力について、より正確な言語を使用していく。

また、AI を活用してみたい、特定のタスクを行う際に AI が役立つ場面と妨げになる場면을繰り返し学習したいという人材も必要である。

好奇心を持ち、適応性があり、他者のことを考え、継続的に学ぶ性質は、今後ますます重要になっていくであろう。

特定のタスク領域、つまり仕事で優れた成果を上げるための一般的な能力を意味する「ジョブコンピテンシー」を備えた AI を開発するのは、まだかなり先のことである。

スライド 34

AI の規制については、活発かつ切迫した議論が続いている。

Deep Mind の共同創設者で、現在は Microsoft AI の CEO を務めるムスタファ・スレイマン氏といった書き手は、AI イノベーションの波によって甚大な被害がもたらされる可能性がある」と警告している。

スレイマン氏の見解は、ダリオ・アモデイ氏とほぼ正反対である。つまり、人々は AI のリスクを認識したくないため、その危険性を過小評価していると言うのだ。

政府、立法者、規制者といった社会の統制システムは、こうした問題にどのように対応するのであるか。今後、やってくる波をどうやって抑え込めば良いのであるか。

そして、記録、レコードキーピング、アーカイビングが社会の統制システムを強化し、有効化し、サポートする場合、最善の支援を行うには、どうしたらよいのであるか。

スライド 35

英国は現在、AI 規制に対して中道的な政策をとっている。個人データの使用を含め、比較的規制の少ない米国のアプローチに従っているわけではなく、また、EU が AI 法で制定した規制インフラのようなものもまだ導入していない。

EU では、AI を提供する側と採用する側にかなりの新しい責任を課しており、法的責任や、EU 市民の基本的な権利への影響を十分に考慮した義務が発生する仕組みを導入している。

AI 法は、これらの法的に強制可能な権利を網羅するために、データ保護規制当局の役割も拡大している。

スライド 36

英国政府機関の AI 導入を支援するため、中央デジタルデータオフィスでは 2024 年 1 月、生成 AI フレームワークを作成した。これによって、TNA での AI の活用に向けた方針を導

入する道が開かれた。

政府の方針は、政府機関、大手ハイテク企業、学者、市民社会、規制当局など、さまざまな関係者との協議を経て策定された。TNA では、現場でこの政策を導入している。

この方針は、明確な規則というよりも、原則に基づくものである。原則を自分たちの状況に当てはめる余地があるので、TNA にとっては望ましいものである。

しかし、ここで私たちは多くのことを自問しなくてはならない。つまり、回答を導き出すために、何度も思考と作業を重ねなければならないということである。

生成 AI とは何かと、生成 AI の限界は、おわかりであろう。

これは、言うは易し、行うは難し、である。

法律に則り、倫理観と責任を持って、生成 AI を活用すること。

この原則は非常に大きな役割を果たしている。

あるコンテキストにおいて、どの法律が AI の活用に適用されるかを把握することは、アーカイブズにとっても決してささいな問題ではない。

そうは言っても、AI の活用が合法か否かを検証するのは、私たちにとってはおそらく最もわかりやすい問題であろう。これは、簡単だからではなく、私たちが最も知識と経験を備えている分野だからである。

一方、アーカイブズの置かれた環境を含め、AI の活用に関して広く受け入れられた倫理規範はまだ確立されていない。

個人と集団の権利および利益の間には、両立しえないものがある。倫理的な考慮事項の一部は、文化特有の要素があり、それは異なるグループ間も同様である。さらに、環境負荷についても考慮しなければならない。

結論

スライド 37

この新たな AI の世界において、アーカイブズやアーキビストはどのような位置付けになるのだろうか。

私たちは、確かに業務に AI を導入し、さまざまな用途に広く活用することはできる。しかし、その方法については慎重に考えなければならない。それは、人間の持つ深い知性、そして AI がもたらす全く種類の異なる知能を理解することを意味する。

何にせよ、多くの分野で多くの誤った認識を持つ人がでてくるだろう。例えば、手を広げすぎて、非現実的なプロジェクトに取り組んだり、あるいは新しい技術に反発して過剰反応する人が出てきたりである。

私たちアーキビストは聡明である。自分たちの状況に合わせて技術を理解し、有効に活用することができる。

情報の保有者として、私たちは良き管理者になる責任がある。特に、当館の歴史的なコレク

ションのデジタル化、AIでの活用、使用方法については、重大な責任を負っている。

また、現代の記録の重要性が認知され、保存されるようにしていくのも私たちの責任である。

つまり、AIを管理するためのアーカイブズの実務を発展させる作業が求められるのである。

社会が整備してきた「統制システム」、つまり政府、立法府、法の支配、規制、規制当局と

いった、私たちの安全を守るためのシステムは、AIの登場によってさまざまな課題にさら

されるようになり、今後押し寄せる波に遅れずについていくのは難しくなっていくだろう。

アーカイブズは説明責任を果たすための基盤の一部である。説明責任のある意思決定を可

能にし、透明性を高め、法の支配をサポートする、そうした記録を保存していく。こうした

統制システムが機能するよう、支援していくのが私たちである。

AIの導入と活用状況をより広く見ていくと、私には、アーキビストの視点、知識、スキル、

感性が必要とされていることは明らかのように思える。

私たちは、AI革命の一員となる必要がある。

そして、AIによって生じる課題に立ち向かわなければならない。

私たちは、課題解決に向かって取り組んでいくことになるだろう。

ご清聴ありがとうございました。