

AIと生徒・教員が共進化する中学校技術科の安全教育教材の提案

Proposal for safety education materials for junior high school technology classes where AI co-evolves with students and teachers

根本 航太¹ 佐藤 里恵² 村越 英樹^{2*}

Kota Nemoto¹ Rie Sato² Hideki Murakoshi^{2*}

¹株式会社ミラプロ MIRAPRO Co., Ltd.

²東京都立産業技術大学院大学 Advanced Institute of Industrial Technology

*Corresponding author: Hideki Murakoshi, hm@aiit.ac.jp

Abstract We propose an AI-assisted, co-evolutionary safety education system for junior high school technology classes, where students engage in practical activities such as woodworking and metalworking that inherently involve safety risks. Current instructional environments face limitations due to large class sizes and the high supervision burden on a single teacher. Moreover, existing textbooks and materials emphasize fundamental safety awareness but do not sufficiently cultivate students' ability to recognize and prevent hazards autonomously. To address these issues, we propose introducing a system that integrates wearable cameras, AI-based hazard detection, and KYT (Kiken Yochi Training) to form a cyclical learning process. During practice sessions, students' actions are recorded and analyzed by AI: the multimodal model CLIP identifies potentially dangerous scenes, while YOLO detects specific hazardous regions within each frame. These detected frames are then converted into individualized learning tasks in which students identify and reason about dangerous points. Their responses are subsequently used to refine both classroom KYT discussions and AI re-training, allowing students and AI to co-evolve. A prototype implementation confirmed that CLIP could recognize hazardous behaviors such as operating a drill while wearing gloves, though prompt engineering remains critical for higher precision. YOLO-World accurately detected common objects (e.g., gloves) but struggled with specialized machinery such as bench drills, indicating the need for domain-specific datasets. This co-evolutionary model promotes both continuous AI improvement and the development of students' active hazard-perception skills, suggesting a novel direction for technology-education safety instruction.

Keywords safety education; technology education; co-evolutionary learning; wearable camera

1 はじめに

中学校技術科は、木材・金属等の材料加工や電気工作などの実習を通して、ものづくりに関する基礎的な知識や技能を育成することを目的としている。しかし、これらの学習活動は工具や機械を使用するため、常に一定の危険を内包している点に特徴がある。独立行政法人日本スポーツ振興センターが公表している学校管理下の災害統計[1]によると、「実習・実験室」で発生した負傷件数 2,087 件のうち、技術・家庭科の授業中に発生したものは 835 件と高い割合を占めている。すなわち、技術科は教育活動の中でも特に事故の発生リスクが高い教科である。

一方で、安全管理の観点からは、学級規模の大きさも重要な課題である。2022 年の法改正[2]により、中学校の 1 学級あたりの生徒数は段階的に 35 人へと引き下げられたが、厚生労働省職業能力開発局の「職業訓練サービスガイドライン」[3]では、実技指導において「受講者 15 人に講師 1 人以上」が望ましいとされている。これと比較すると、技術科の授業は依然として 1 人の教員に過大な安全管理負担がかかっており、生徒一人ひとりの作業状況を常時監督することは困難である。すなわち、「クラス定員の多さ」が安全指導の限界を生む原因の一つとなっていると考えられる。

さらに、教育現場の調査[4]によれば、事故やヒヤリハットの多くは工具や機械の扱い時に集中しており、特に電動工具の使用場面に危険が偏在していることが指摘されている。また、安全衛生に関する授業時間が設けられていない学校も少なくなく、注意喚起の掲示や安全マニュアルの整備が不十分な場合も多い。作業時の服装についても、安全靴や作業服を着用せずに制服やスリッパで実習を行う例があり、環境整備・服装面の不備も事故の一因となっている。このように、道具・環境・安全教育の三要素が複合的に絡み合い、安全指導の実効性を損ねている現状がある。

川路・谷田(2020)[5]は、現行教科書における安全教育の内容を分析し、「安全能力」の育成がどのように意図されているかを明らかにした。その結果、授業中の工具取扱い、家庭内での安全確保、社会的安全意識の涵養といった基礎的な記述が中心であり、危険情報の整理・行動ミス防止・不安全行動の自制など、初歩的な安全意

識の定着に重きが置かれていることが示された。一方で、経験に基づく危険察知、原因分析、再発防止といった高度で実践的な安全能力は十分に扱われていない。すなわち、教科書は「安全行動の基礎」は育成できるが、「現場での危険判断力」や「主体的な再発防止力」までは育成しきれていないという限界を持つ。

このように、教科書を中心とした従来の安全教育には「現場対応力の欠如」という構造的な課題があり、生徒が自ら危険を察知・判断する能力を育成するには不十分である。したがって、今後の安全教育には、学習者自身が危険を「見て・考えて・判断する」プロセスを体験的に学べる教材の導入が求められる。

文部科学省の「中学校技術・家庭科(技術分野)事例集」[6]では、安全指導を含む授業実践の工夫が紹介されているものの、特定単元に即した断片的な事例が多く、体系的な安全教育カリキュラムの構築には至っていない。また、「技術・家庭科安全ハンドブック」[7]など、教員向けの安全指導資料も整備されているが、これらは主に指導者視点からの安全管理を目的としており、生徒が自ら危険を見つけ、考える学習機会は十分に設計されていない。

教育実践研究の中では、江口(2012)[8]が中学生向けの KYT (危険予知トレーニング)教材を JavaApplet で開発し、グループ学習の制約を超えた個別学習型の試みを行っている。この教材は、即時フィードバック機能によって生徒の危険予測力を高める効果を示したが、提示される危険場面はあらかじめ用意された静止画像であり、生徒自身の実体験や学校固有の環境を反映することは難しかった。また、教材の更新や個別化には教員の手作業が必要であり、現場への普及には課題が残る。

さらに、技術科教育における安全指導を ICT 活用の観点から扱った事例は限定的である。山下・田中・谷田(2024)[9]は、技術科教員の教材提示手法を分析し、ICT 機器が整備されているにもかかわらず、教科書や実物提示といった従来型の指導法が依然として主流であることを明らかにしている。この傾向は ICT 機器を使用した学習指導方法に対する教員の熟練度が影響を与えていると推察している。

一方で、教材開発技術の分野では、近年、AI や自然言語処理を用いた問題自動生成や教材自動生成の研究が進展しているが

(例: 石田ら, 2025[10])、これらの多くは理数系教科や語学など、知識理解を対象とした学習支援が中心である。VR を利用した体験型の教材開発(庄 司ら, 2022[11])や安全教育教材(武田ら, 2020[12])も進んでいる。模擬的な体験を提供する点では有効であるものの、現実の作業現場や生徒個々のデータに基づいて危険を可視化し、学習へ還元する仕組みは構築されていない。

関(2020)[13]は医療現場で手術のトレーニングをARで行うシステムを開発した。熟練医の手技映像を学習者の視野に重ね合わせて模倣学習を促すシャドウイング型の訓練支援を主目的としている。そのため、主に動作の再現性や立体的な手技理解の支援に重点が置かれており、危険な操作の自動検出や誤った手技の判定といった安全面のフィードバック機能は備えていない。

以上のように、中学校技術科における安全教育の現状には、複数の課題が複合的に存在する。第一に、授業環境そのものが工具や機械を扱う高リスクな性質を持ちながら、一人の教員が多数の生徒を監督する構造的制約が存在し、作業を細かく把握することが難しい。第二に、教科書や既存教材は安全意識の基本教育には有効であるものの、生徒が自ら危険を発見・判断し、再発防止策を考えるような実践的安全能力を育成する設計になっていない。第三に、ICT や AI 技術の進歩が進む一方で、その利用は教員の個々の能力に左右され、実際の現場の作業映像や生徒自身の行動データを教材化する仕組みはほとんど整っていないといえる。

既存研究では現場の多様な環境や生徒の個人差に応じて動的に危険を可視化し、学習者自身の体験に還元する仕組みは未だ確立されていない。

したがって、本研究では、中学校技術科における安全教育の課題を解決するために、AI による危険行動検出と学習者主体の危険予知活動を組み合わせた「生徒主体型・共進化型教材」の設計を目的とする。

本稿第 2 章では、現場の危険行動データを基にした AI・ウェアラブル教材が提供するサービスの構想について述べる。第 3 章では危険行動を検出するウェアラブルデバイス、および教材開発について説明する。第 4 章では試作したウェアラブルデバイスの危険行動検出評価について記載する。第 5 章は本稿のまとめである。

2 サービス構想

本サービスの全体像を図 1 に示す。ウェアラブルカメラ、AI 解析、危険予知訓練(KYT)を組み合わせ、記録・抽出・課題化・改善の循環を通じて AI と生徒が共に成長する共進化型の安全教育モデルとして設計されている。以下では、図 1 に示した各フェーズごとにサービスの詳細を記載する。

記録

実習時、生徒はウェアラブルカメラを装着し、作業過程を継続的に記録する。これにより、従来は教員の目視に依存していた観察を客観的データとして蓄積でき、個々の生徒の作業状況を時系列で追跡可能となる。映像は授業後の解析に使用される。記録段階で取得する映像によって、後段の危険推定と教材生成の効率化に資する。

また、授業中に危険行動が検知された場合は、抽出フェーズから

アラートを受け取ることで、アラームを鳴らして危険行動の発生を知らせることで教員を補助する。

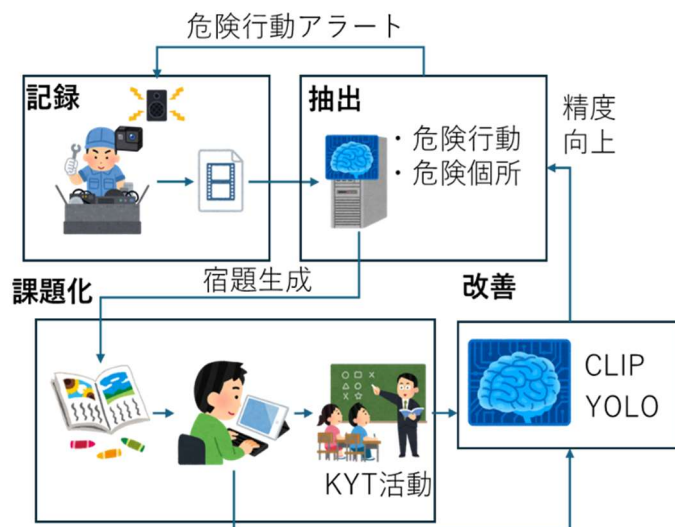


図 1 サービスの全体像

抽出

映像に対して AI 解析を行い、危険関連場面と具体的な危険箇所を段階的に抽出する。まず CLIP により「手袋を着用したまま卓上ボール盤を操作する」等の自然言語プロンプトとの類似度計算を用いて潜在的に危険度の高いフレーム群を自動選別する。この作業は作業中に行われ、危険行動が抽出されたときはアラートを発する。

危険行動だと判定されたフレームは YOLO により物体検知が行われ、フレーム内の危険部位の位置を特定する。

こうして得られたデータは、後段の「課題化」に使用され、AI と生徒・教員が共に成長する「共進化」の基盤となる。

課題化

抽出された静止画像は宿題として生徒に配信される。生徒は提示画像上で危険箇所を指示し、その理由を考察して回答を提出することで、「見て・考えて・判断する」プロセスを反復的に体験する。

さらに、授業内では KYT(危険予知トレーニング)として、宿題で生徒が誤って回答した画像を用い、グループ討議によって危険性・想定事故・対策を検討する。その後、各グループは発表を通じて知見を共有し、その内容は AI にフィードバックされる。

本設計にはナッジの要素を取り入れており、安全行動が増えると宿題量が減少する仕組みを通して、生徒は安全行動の内在的報酬を経験的に学ぶ。教材は単なる知識伝達の道具ではなく、生徒の行動変容を促す仕組みとして機能する。このサイクルを通じて、生徒の学びと AI の検知能力の双方が向上する共進化が実現する。

改善

KYT 活動で得られた生徒の討議内容や新たに発見された危険事例は、AI の学習データとして活用される。具体的には、討議の中で挙げた新しい危険行動を自然言語プロンプトとして CLIP に登録し、それに対応する新しい物体ラベルを YOLO に定義する。これにより、AI はこれまで検出できなかった危険要素を新たに識別可能となる。

次に、登録したラベルを用いて YOLO-World で自動抽出を行い、抽出結果を宿題として生徒に提示する。生徒は提示された画像の正誤を判定し、その回答内容が YOLO の再学習用データとして蓄積される。このようにして、学習者の経験をもとに AI の検出精度が段階的に高まる共進化的改善サイクルが形成される。

3 危険行動検出と教材開発

本章では、本サービスを支える AI 技術として、CLIP を用いた危険場面抽出と YOLO を用いた危険箇所特定の詳細と、それらを使用した教材作成の仕組みについて述べる。

AI の処理フロー

図 2 に、本研究での 2 種類の AI を組み合わせた処理フローを示す。まず、生徒がウェアラブルカメラで撮影した実習動画から、CLIP によって危険が含まれる場면을抽出し、続いて YOLO によって危険箇所を特定する。

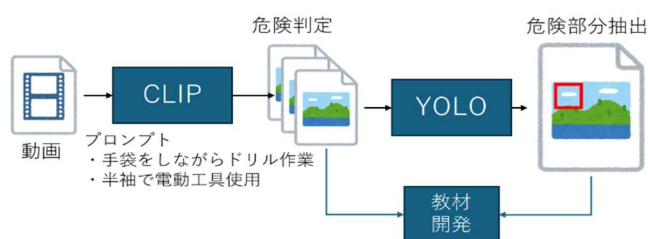


図 2 AI の処理フロー

危険判定

本研究における危険判定は、OpenAI が開発したマルチモーダルモデル CLIP (Contrastive Language-Image Pretraining) [14] を基盤とする。CLIP は、膨大な画像と言語の組み合わせを事前学習することで、テキストと画像を共通のベクトル空間に写像し、その類似度を計算できる特性を持つ。任意の文章 (例: 「手袋をしたままドリルを使用している」「腕時計をしたままドリルを使用している」) を入力すると、その内容と関連性の高い画像を自動的に判別できる。

この特性を活用し、危険行為を自然言語で定義し、映像中の各フレームをテキストと比較することで危険判定を行う。

具体的な処理の流れは以下の通りである。

①危険行動のテキスト定義

教育現場で想定される代表的な危険行為を自然言語で記述するとともに、逆の意味となる文章、すなわち安全であると認識される文章を記述する。

例えば、「保護メガネを着用せずに電動工具を使用している」の対として、「保護メガネを着用せずに電動工具を使用していない」などの組み合わせを複数記述する。危険な行為と安全な行為の類似度の平均値を取り、その差を危険行動のスコアとする。

②閾値の設定

CLIP はプロンプト内容によってスコアが変動するため、あらかじめ「安全な作業映像」を基準データとして撮影し、①で定義した危険行動のスコアから危険判定の閾値を定義する。

③テキストとの類似度計算

ウェアラブルカメラで撮影した実習映像をフレーム単位に分割し、①で定義した危険行動のスコアを算出する。閾値を超えた場合、そ

のフレームを「危険」と判定する。

④危険場面の抽出

③で危険と判定されたフレームを自動的に抽出し、教材 (宿題や KYT 活動用の資料) として登録する。

危険行動を認識する AI モデルは他にも存在するが、本研究で CLIP を採用した理由は主に 2 点である。

1 点目は軽量な点である。CLIP は単一フレームごとの画像認識を行うため、動画全体を解析する行動認識モデルに比べて計算負荷が軽い。そのため、ウェアラブルカメラから送信された映像を AI サーバで逐次解析し、教員が目を離れた場面でも危険を実時間に近い検知が可能であり、危険な行動があった際の警告が可能となる。

2 点目は柔軟な拡張性である。危険行為を自然言語で追加できるため、KYT 活動で生徒が新たに発見した危険を、そのままテキストとして AI に入力するだけで検出対象を拡張できる。すなわち、生徒は AI が提示する危険事例を教材として安全意識を高め、同時に AI は生徒の発見を学習資源として検出精度を高める。この仕組みは「人と AI が共に成長する共進化型の学習モデル」という本研究の理念に適合するものである。

教材開発

CLIP は映像全体を対象に「この場面に危険が含まれているか」を判定できるが、具体的にどの部分が危険であることを示す機能は持たない。そのため、CLIP が危険と推定した画像を入力として、危険箇所を領域として特定する AI モデルを導入する必要がある。本研究では物体検出アルゴリズムの一つである YOLO (You Only Look Once) [15] を採用した。

YOLO は、入力画像を一度の畳み込み処理で解析し、画像内に存在する複数の物体を同時に検出・分類することが可能なモデルである。従来の物体検出手法 (R-CNN 系など) が「候補領域の抽出 → 各領域の分類」という段階的処理を必要としたのに対し、YOLO は画像をグリッドに分割し、各セルに対して物体の存在確率と位置 (バウンディングボックス座標) を直接予測する。これにより、リアルタイム処理が可能な高速性と、位置推定における一貫性を両立している。

本研究における YOLO の役割は、CLIP によって「危険を含む」と判定された画像から、さらに危険とみなされる具体的な部位を抽出し、矩形領域として抽出することである。例えば、手袋を着用したまま工具を操作する手、回転工具に近づいた腕時計、あるいは保護具を着用していない顔周辺といった箇所を対象とする。さらに、CLIP に登録した危険行動のテキスト定義と関連付ける形で、それぞれの危険行動に対応する特徴的な物体を YOLO に登録しておく。これにより、生徒が画像上でタッチした位置が、定義された危険行動に該当するかどうかを判定できる。

YOLO はその高い検出精度と汎用性から、多くの派生モデルが開発されている。その中でも YOLO-World [16] は、アノテーションを必要とせず、自然言語で記述したラベルに基づいて物体を検出できる点に特徴がある。

ただし、YOLO-World は教師データを用いた従来のモデルに比べて検出精度が低下する場合がある。そこで本研究では、YOLO-World が抽出した結果をアノテーション支援データとして活用する。具体的には、新たに登録された物体に対して、蓄積された画像群から YOLO-World で検出を行い、その結果が正しいか否かを生徒が

二択で判定するユーザインタフェース(UI)を実装する。これにより、バウンディングボックスを手作業で作成せずとも、簡易的なアノテーションが可能となる。さらに、このプロセス自体を「宿題」として学習サイクルに組み込むことで、生徒がタブレット上で判定作業を行うたびにアノテーションデータが蓄積され、YOLO モデルの再学習による精度向上が実現する。

4 危険行動抽出評価

図 2 に示した危険行動抽出装置を試作し、CLIP による危険行動の抽出精度と YOLO による危険部位の抽出精度の検証を行った。本検証の目的は、CLIP モデルによって「手袋を着用した卓上ボール盤操作」を自動的に識別できるかを検証することである。

CLIP による危険行動抽出

今回使用する映像は、都立産業技術大学院大学の夢工房内で卓上ボール盤を使用する場面を撮影した。図 3 左に示す「手袋を着用せずに卓上ボール盤を使用する」正しいシーンと図 3 右に示す「手袋を着用して使用する」危険なシーンの 2 パターンの行動を撮影した。



正しいシーン

危険なシーン

図 3 正しいシーンと危険なシーンの比較

危険行動検出実験には、撮影した動画を 16 本に分割し、正しい行動の動画 10 本と手袋をしている危険な動画 6 本を検証に使用した。

CLIP のモデルは「CLIP-RN50x64」を使用し、プロンプトは図 4 に示す。「手袋をしながらドリルを操作している」危険行動を説明する文章を 9 件、危険行動に該当しない文章を危険行動を説明する文章に対になるように 10 件記載した。

危険行動および安全行動を表す各プロンプト文を、図 5 に示す 6 種類のテンプレート文に挿入することで、多様な文章表現を生成した。これにより、特定の文体や語句に依存した認識の偏り(バイアス)を抑制し、より安定した判定が得られるようにした。

作成した各テンプレート文は CLIP モデルに入力され、テキスト側のベクトルとしてエンコード・正規化を行った。同様に、映像フレームを CLIP で画像ベクトルに変換し、両者のコサイン類似度を算出した。最終的には、危険行動(positive)プロンプト群と安全行動(negative)プロンプト群の類似度の差(マージン値)をスコアとして用いた。

閾値の設定は、撮影した映像から不安全なシーン、安全なシーンの画像を 7174 枚抽出してスコアを算出した。このスコアを利用して、不安全な動画を 1 度でも不安全だと判定できる値(0.019)を閾値とした。将来的にはこの手順を自動で行うようなアルゴリズムを提案したいと考えている。スコアが閾値を 3 フレーム以上連続して超過した場合に「危険行動」と判定した。これは一部フレームでスコアが一時的

に高騰する(スパイクする)現象を抑制し、過検出を防止するための仕組みである。

```
dict(
  name="gloves_on_drillpress",
  #Positive: 「手袋をしながらドリルを操作している」文面
  pos=[
    "Operating a benchtop drill press while wearing work gloves",
    "Using a bench drill with gloves on; hand feeding the rotating spindle",
    "A gloved hand lowering the quill of a drill press during drilling",
    "Close-up of gloves near the rotating drill press chuck while drilling",
    "Gloved operator pressing the lever of a benchtop drill press",
    "Work gloves visible while drilling a workpiece on a drill press",
    "Gloved hands holding the workpiece under a spinning drill press bit",
    "Column drill (pillar drill) in use; operator wearing gloves",
    "Pedestal drill actively drilling; operator has gloves on"
  ],
  # Negative: 危険行動には該当しない文面 (素手・停止中等)
  neg=[
    "Operating a drill press with bare hands and without gloves",
    "Using a bench drill safely with no gloves; bare skin hands visible",
    "Drill press in use; operator not wearing gloves",
    "Using a handheld cordless drill with gloves (not a drill press)",
    "Operating a lathe with gloves (not a drill press)",
    "Using an angle grinder while wearing gloves (not a drill press)",
    "Operating a milling machine with gloves (not a drill press)",
    "Bench grinder with gloves on (not drilling)",
    "Drilling on a benchtop drill press with bare hands, no gloves",
    "Drilling on a benchtop drill press with bare hands"
  ]
)
```

図 4 実験に使用したプロンプト

```
"Photo of {} in a workshop.",
"Video frame of {} at a school tech lab.",
"A person {} in a factory workshop.",
"Close-up photo of {}.",
"High-speed video still of {}.",
"Safety training image showing {}."
```

図 5 実験に使用したテンプレート文

各画像を 100 回ずつ評価したところ、危険行動時の画像 6 本は 100% 検知できている。しかしながら、安全行動時の画像 10 本は 50% 検知にとどまっている。安全行動時に危険行動と誤検知している原因として、手袋とドリル部が同一フレーム内に十分映っていない映像でスコアが低く出てしまい、それに閾値を合わせてしまった結果として誤検知が増えてしまった。スコアが低い 1 本を見逃す閾値を設定すると、誤検知は 0 本になることがスコアの分布からわかっている。安全上、見逃しを減らすことが重要になってくる。したがって、危険行動の見逃しを最小限に抑えつつ、誤検知を低減するための閾値設定が、今後の精度向上における重要な課題であると考えられる。

使用した CLIP-RN50x64 は高精度な ResNet ベースモデルであり、危険行動と安全行動のスコア差を比較的明瞭に捉えることができた。一方で、推論負荷が大きく、教育現場でのリアルタイム運用には軽量モデル(例: ViT-B/32, RN50x4)との比較や最適化が必要である。さらに、CLIP モデルの特性上、入力テキストが英語である点は教育応用上の制約となる。学習者が日本語で危険シーンを記述し、それを AI が解釈できるようにするためには、日本語 CLIP モデルの導入や日英対応プロンプトの検討が求められる。

なお、CLIP のプロンプト設計では、テンプレートや文章表現の違

いによってスコアが大きく変動することが確認された。このことは、AI にとって「どのような表現で危険を説明するか」が識別精度に直結することを意味する。教育応用の観点からは、生徒が自ら危険行動の具体的な場面を想起し、それを言語化して AI に提示する過程自体が安全意識の向上につながると考えられる。今後は、プロンプト設計による検出率の違いを可視化し、生徒がその効果を体験的に理解できる仕組みを検討していく。

YOLO による判定

今回は「手袋をしながら卓上ボール盤を使用」するシーンの中で危険箇所の検出をした「手袋」と「ドリル」のラベルを追加し、YOLO-World にて抽出を行った。

「手袋」は図 6 で示すようにかなり正確に検出できている。一方で、素手を「手袋」として検出している部分もあるので学習用画像とする前に判別する必要がある。

一方で「ドリル」(卓上ボール盤)は様々なラベルを試してみたが、卓上ボール盤自身をドリルとして認識することはほとんどなく、別の物体をドリルとして検出してしまった。

卓上ボール盤のような専門的な物体は学習データに含まれていないため YOLO-World では判別できないためと推察できる。



図 6 YOLO を使った判定結果

5 おわりに

本研究では、中学校技術科における安全教育の課題を踏まえ、ウェアラブルカメラと AI を活用した共進化型安全教育システムの構想を提案した。CLIP を用いて実習映像から危険行動を抽出し、YOLO によって具体的な危険箇所を特定することで、生徒が「危険を見て・考え・判断する」学習を日常的に行える仕組みを示した。

試作では、CLIP および YOLO-World を用いた危険判定の有効性を検証した。その結果、一定の精度で危険行動を識別できることが確認された。CLIP による判定では、プロンプト設計やテキスト表現の工夫によって精度向上の余地があると考えられる。

一方、YOLO-World による物体検出では、「グローブ」のような一般的な対象は比較的高い精度で認識できたのに対し、「卓上ボール盤(ドリル)」のような専門的機器については検出が不安定であり、誤認識も多く見られた。これは、事前学習モデルの学習データに産業用機器が十分含まれていないことが原因であると推察される。

これらの結果から、一般物体検出モデルをそのまま安全教育の教材生成に適用するには限界があることが示唆された。したがって、今後は手動アノテーションを通じて教師データを構築する仕組みを導入する必要がある。また、CLIP による危険文脈の判定と YOLO による物体検出を組み合わせて、より精度の高い危険箇所特定が可能になると考えられる。

今後は、アプリケーションの開発を進めると共に実際の学校授

業への導入に先立ち、教員および生徒を対象としたワークショップ形式の検証実験を実施する予定である。ワークショップでは、模擬的な実習環境を設定し、参加者にウェアラブルカメラを装着して作業を行ってもらう。その映像を CLIP および YOLO で解析し、AI の危険検知結果と人間の危険認知との対応関係を比較・検証する。また、参加者自身が AI の判定結果をもとに危険箇所を考察する体験を通して、AI との協働が安全意識の変容に与える影響を観察する。これにより、本システムの教育的有効性と改良点を明らかにすることを目指す。

さらに、CLIP のプロンプト設計や YOLO-World による自然言語ベースの危険物体登録機能を改良し、参加者が自ら新たな危険行動を定義・追加できる仕組みを検討する。プロンプト作成や危険シーンの言語化を通じて、学習者が安全を自ら構築する主体的学びを促進できる可能性がある。

今後の課題としては、ワークショップ環境と実際の学校環境との乖離、映像データの扱いに関する倫理的配慮、AI モデルの一般化性能などが挙げられる。これらの課題を段階的に解決しながら、将来的には中学校現場における安全教育の補助教材としての実装を目指す。AI と生徒が共に成長する本システムを通じて、技術教育における新しい学びの形を探索していきたい。

参考文献

1. 独立行政法人日本スポーツ振興センター. 学校等の管理下の災害[令和 6 年版]. 2025. Available from: <https://www.jpnsport.go.jp/anzen/kankobutuichiran/tabid/3053/Default.aspx>
2. 文部科学省. 公立の義務教育諸学校等の教育職員の給与等に関する特別措置法等の一部を改正する法律案(概要). 2025. Available from: https://www.mext.go.jp/b_menu/houan/an/detail/mext_03106.html
3. 厚生労働省. 職業訓練サービスに関する指針. 2025. Available from: https://www.mhlw.go.jp/file/06-Seisakujouhou-11800000-Shokugyounouryokukaihatsukyoku/02_5.pdf
4. 磯部 征導, 宮川 秀俊, 村松 庄太郎. 愛知県内の中学校技術科教育における安全管理と安全指導の現状と課題. 日本産業技術教育学会誌. 2017;59(1):1-8.
5. 川路 智 治, 谷田 親彦. 安全能力の構造に基づいた技術科教科書の分析. 日本教科教育学会誌. 2020;43(3):1-9. Available from: https://www.jstage.jst.go.jp/article/jcrdajp/43/3/43_1/_pdf/-char/ja
6. 文部科学省. 中学校技術・家庭科(技術分野)事例集. 文部科学省. 2025. Available from: https://www.mext.go.jp/a_menu/shotou/new-cs/senseiounen/mext_02685.html
7. 矢澤 巖弥, 田中 伸英, 山口 三枝子, 濱 晴奈, 布川 広, 野田 まなみ. 技術・家庭科における安全指導の工夫—すぐに活用できる技術・家庭科安全ハンドブックの作成—. 川崎市教育委員会. 2025. Available from: https://kawasaki-edu.jp/index.cfm/7%2C226%2Cc%2Chtml/226/25-149-152.pdf?utm_source=chatgpt.com
8. 江口 啓, 渡邊 肇也, 小林 健太, 金原 恭. 中学生のための KYT シートを用いた安全教育教材の開発. 日本産業技術教育学会誌. 2012;54(4):205-212.
9. 山下 大吾, 田中 愛也, 谷田 親彦. 中学校技術科における ICT 機器等を活用した学習指導の実態調査. 日本教育工学会論文誌. 2024;48(Suppl.),113-116. Available from: https://www.jstage.jst.go.jp/article/jjet/48/Suppl./48_S48060/_pdf/-char/ja
10. 石田 崇, 雲居 玄道, 小林 学, 平澤 茂一. 生成 AI を用いた統計学の学習用練習問題自動生成システムの試作. 日本教育工学会論文誌. 2025;49(2):341-354. Available from: https://www.jstage.jst.go.jp/article/jjet/49/2/49_48068/_pdf/-char/ja
11. 庄司 真史, 小林 佑介, 河合 研志, 佐藤 友彦. 視点の水平移動可能な BYOD 型地学 VR 巡検教材の開発. 地学教育. 第 74 巻第 1 号, 13-30, 2022. Available from: <https://www.jstage.jst.go.jp/article/chigakukyoiku/74/1/74>

- [_13/_pdf/-char/ja](#)
12. 武田 昂大, 湯原 聖也, 木原 拓己, 栗飯原 萌, 古市 昌一. 指差呼称による玉掛け作業の安全教育を目的とした VR システムの開発. 日本デジタルゲーム学会 年次大会 予稿集. 第 10 回 年次大会, 183-186, 2020. Available from: https://www.jstage.jst.go.jp/article/digraiproc/10/0/10_183/_article/-char/ja/
 13. 関 拓哉. Microsurgery 用 AR トレーニングシステムの開発. ライフサポート. 32 巻 1 号, 17, 2020. Available from: https://www.jstage.jst.go.jp/article/lifesupport/32/1/32_17/_pdf/-char/ja
 14. Radford A, Kim JW, Hallacy C, Ramesh A, Goh G, Agarwal S, Sastry G, Aspell A, Mishkin P, Clark J, Krueger G, Sutskever I. Learning transferable visual models from natural language supervision. arXiv:2103.00020 [Preprint]. 2021 [cited 2025 Aug 16]. Available from: <https://arxiv.org/abs/2103.00020>
 15. Redmon J, Divvala S, Girshick R, Farhadi A. You only look once: Unified, real-time object detection. arXiv:1506.02640 [Preprint]. 2016 [cited 2025 Aug 16]. Available from: <https://arxiv.org/abs/1506.02640>
 16. Cheng T, Song L, Ge Y, Liu W, Wang X, Shan Y. YOLO-World: Real-time open-vocabulary object detection. arXiv:2401.17270v1 [Preprint]. 2024 [cited 2025 Aug 16]. Available from: <https://arxiv.org/abs/2401.17270v1>