

一人称視点映像とフィジカルAIを融合した行動分析システム

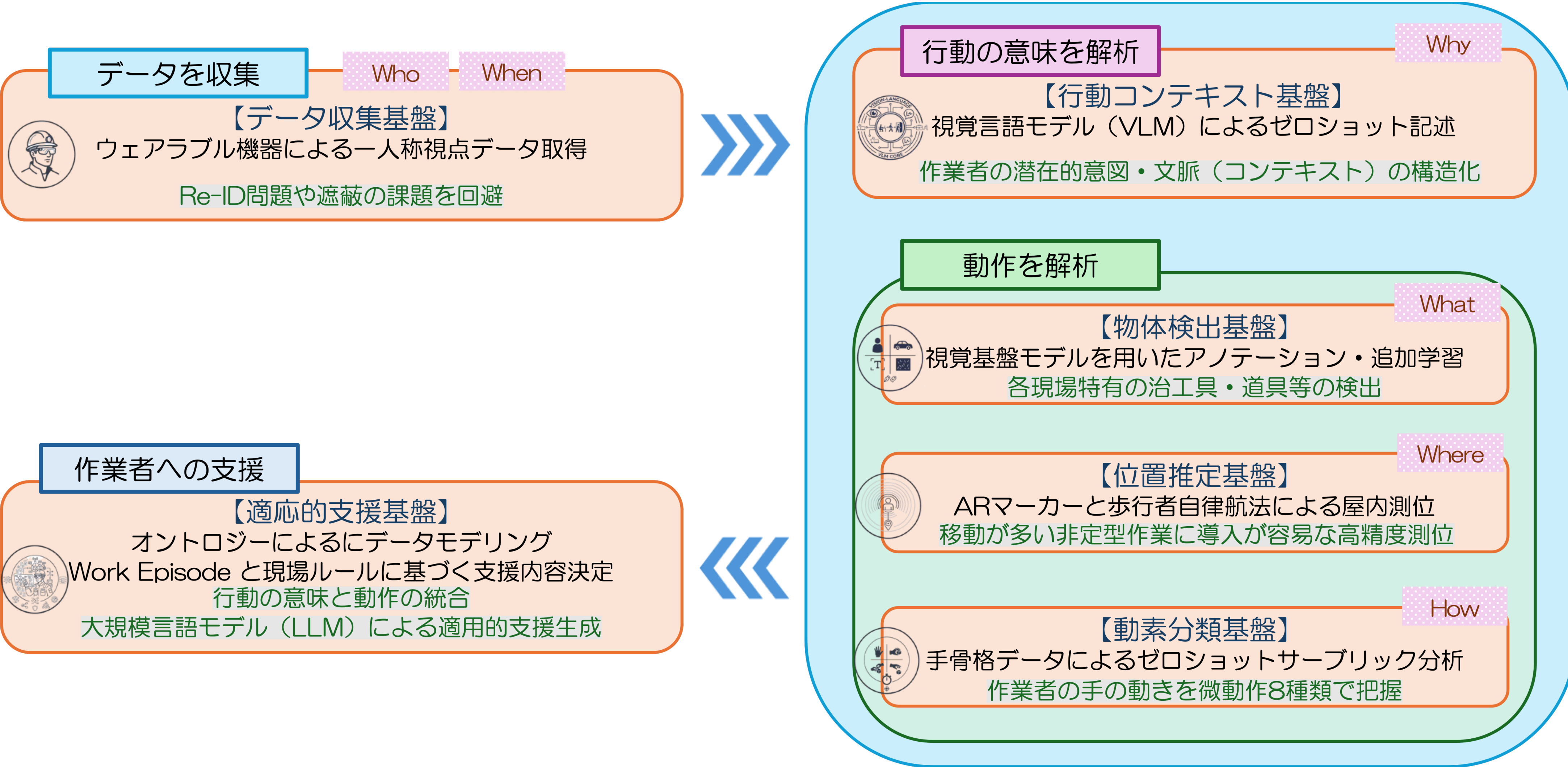
技術概要

ウェアラブル機器の**一人称視点映像**と**モーションデータ**を融合させた、独自のフィジカルAI技術です。作業者の行動の意図や動作をリアルタイムに理解し、現場改善を可能とするのみならず、作業者への最適なサポートを実現します。

3つのポイント

- ①圧倒的な開発効率
視覚基盤モデルを活用することで、現場に合わせた専用のAIモデルを効率的に素早く構築できます。
- ②非定型作業の行動分析
現場のマクロとミクロの動きを把握し、リアルタイムに紐づけることで、**非定型作業**を含み行動分析をすることが可能です。
- ③作業者に合わせた即時サポート
独自のデータモデリングにより行動の意図を深く理解し、作業者一人ひとりに応じた**適応的支援**を可能とします。

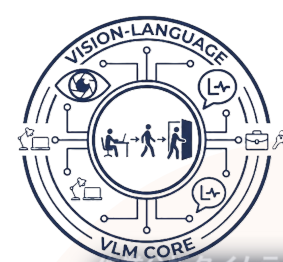
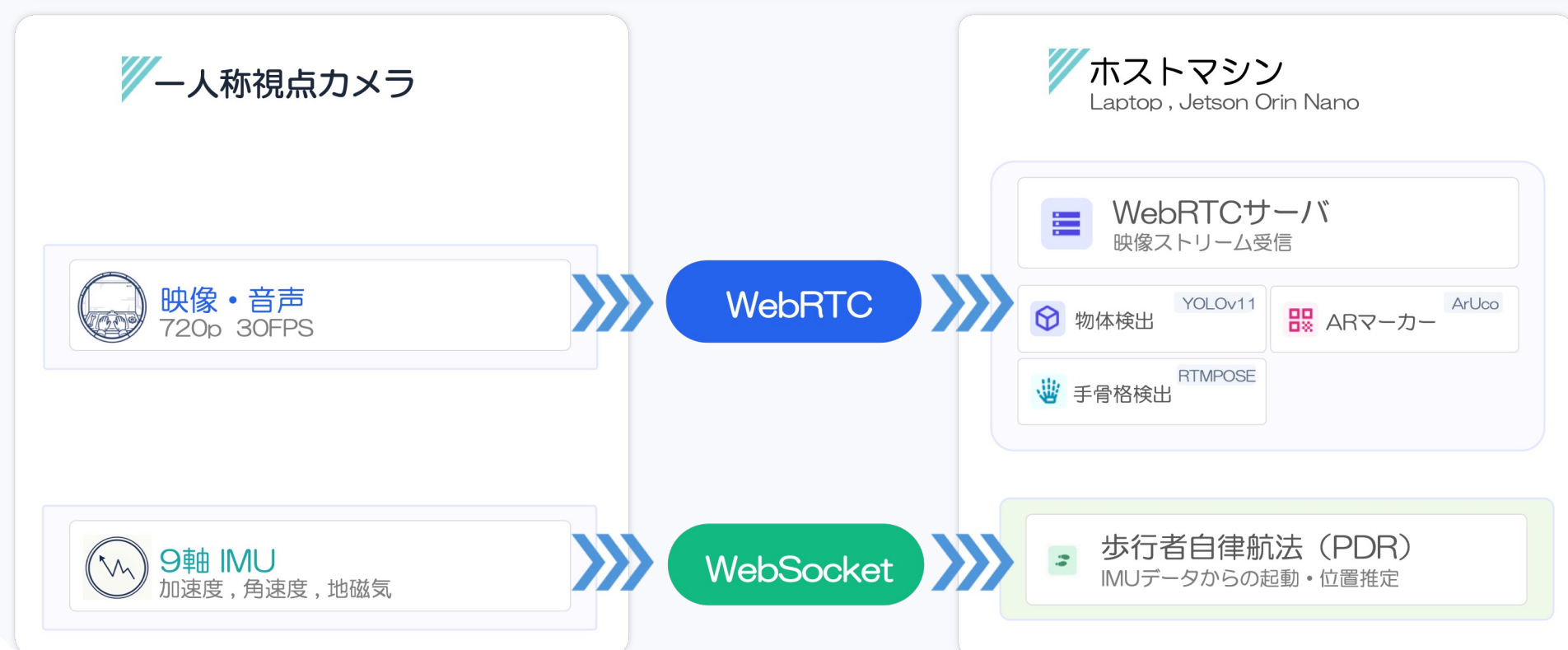
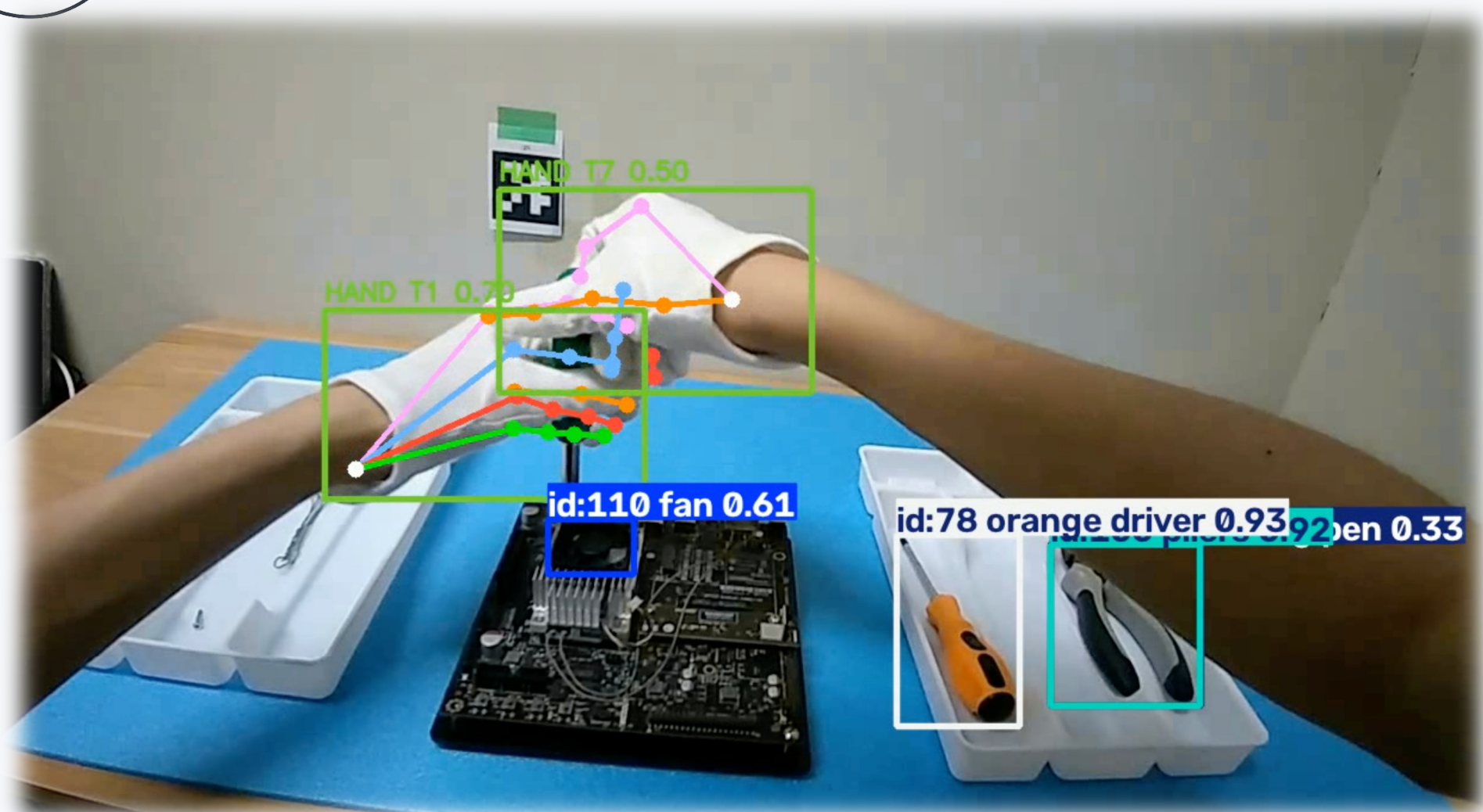
行動分析システムの概要



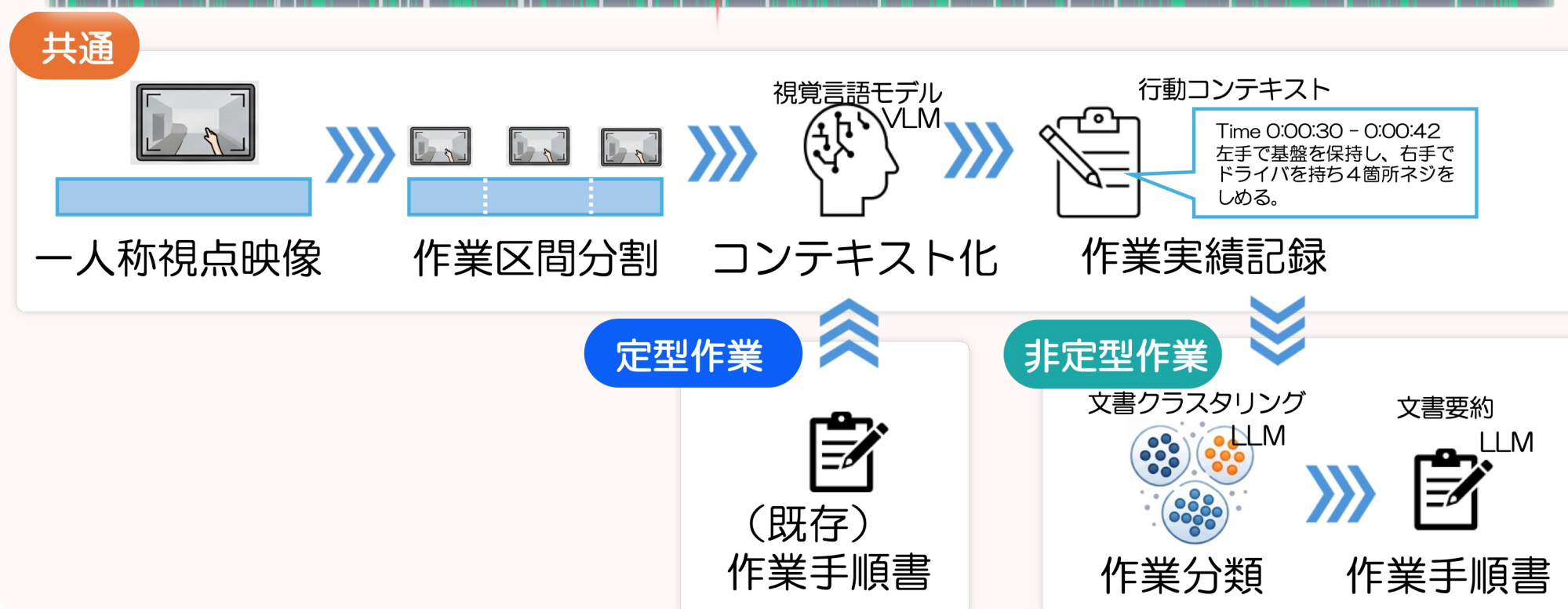
行動分析システムの各基盤



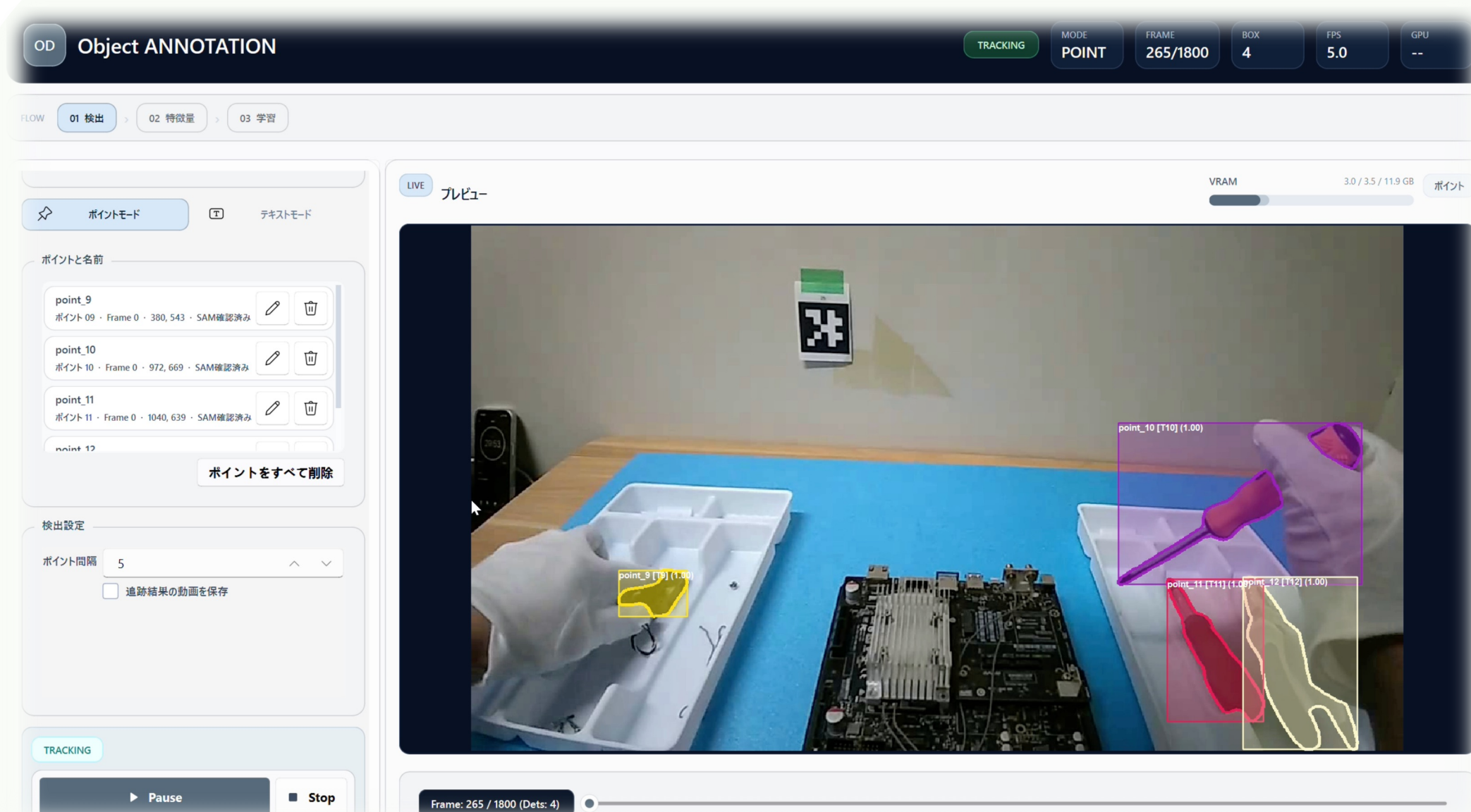
データ収集基盤



行動コンテキスト基盤



物体検出基盤



非定型作業 GroudingDINO

- 治工具がメイン
- 検出したい物体の総数が不明瞭

ドライバー

テキストプロンプトによる物体検出

定型作業 SAM3

- 治工具 + 組立部品パーツ
- 検出したい物体があらかじめ定義

物体マスクを生成

物体候補の精査

候補群

正例を選択

コサイン類似度比較

DINOv3

特徴ベクトル

カスタムモデル作成

データ拡張

軽量物体検出モデル

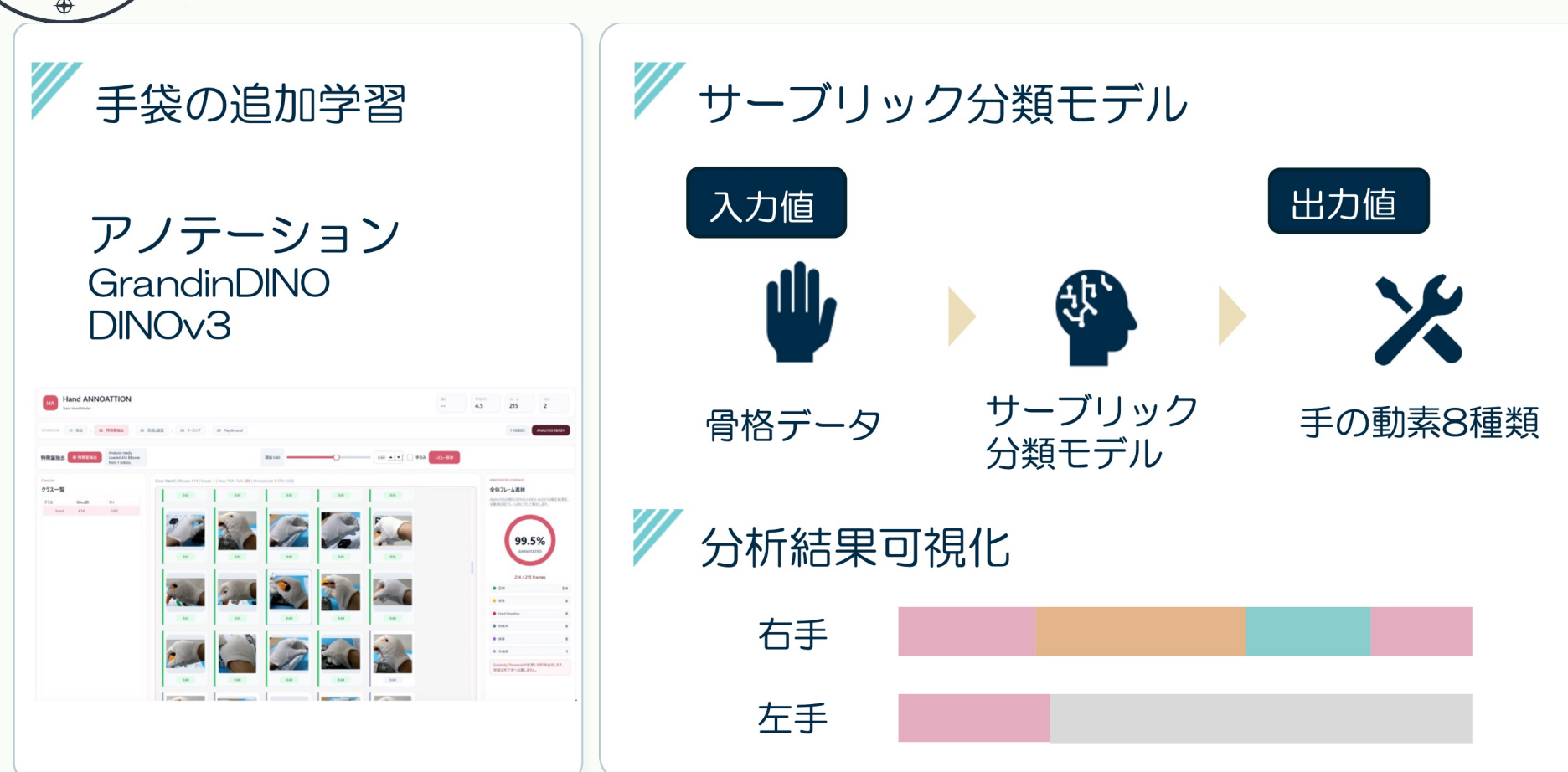
リアルタイム推論



位置推定基盤



動素分類基盤



適応的支援基盤

