

奥行知覚ビジョン

デプスカメラ（2 セット） グローバルシャッターと広い視野「広角」。

最小検知深度は約 0.11m です。

最大 1280 x 720 の深度解像度。

視覚走行距離計カメラ（1 セット）：

高度に最適化された V-SLAM、1%未満の閉ループオフセット、6ms 未満のジェスチャーモーションとモーション反射の間の遅延。魚眼レンズイメージャーと $163\pm5^\circ$ の半球に近い視野を組み合わせることで、高速移動による高速追跡が可能になります。

深度ビジョン-3D マップのリアルタイム作成およびナビゲーション計画：

3D 環境構築：

運動の過程で、ロボットはカメラを使用して環境の色と深度の情報を取得し、特定の視覚アルゴリズムを使用してオブジェクトの 3D 空間情報を再構築します。

確率マップ：

Octomap（確率マップ）は、ロボットが動くと同時にロボットの周囲を検出して障害物データを提供するカメラを使用して構築されます。

動的障害物知覚：ロボットが動的障害物に遭遇すると、ロボットは現在のマップデータを特定の範囲内で更新し、動的障害物がマップ上に残した*移動アーティファクト*を破棄します。

グローバルポジショニング：マップを作成中、グローバルおよびローカルのリアルタイムポジショニング機能を利用することができます。マップはカメラの視点にリアルタイムで追従し、リアルタイムのズームイン、ズームアウト、移動、任意の回転をサポートします。

ループ検出：ロボットは、広範囲のフィールドで高いループバック精度を維持し、特定の範囲内で高い位置決め精度を維持し、ドリフトまたは損失があっても、特定の振動振幅内で安定性を維持できます。

人間の姿勢認識追跡と顔認識

体位認識

カラーカメラは、ディープラーニングモデル(深層学習)に従って人の特定の姿勢を識別し、ヒューマンマシンインタラクションを行うことができます。ロボットはさまざまな体の姿勢に応じて対応する動作を行うことができます。

人間の骨格の知覚

ロボットは、視点からの色情報に応じた人体の2次元骨格情報を解析・計算し、さらに被写界深度を利用して特定のキャラクターの3次元骨格情報や動作情報を解析・計算することができます。

対象者追跡

シーンに複数の人物がいる場合、ロボットに誰をロックするかを特定の姿勢をとることで指示することができます（たとえば、左手を上げる）。その後、ロボットはターゲットの動きに追従します。