

Multi-Scale Deep Reinforcement Learning for Real-Time 3D-Landmark Detection in CT Scans

Florin C. Ghesu^{1,2}, Bogdan Georgescu², Yefeng Zheng², Sasa Grbic², Andreas Maier¹, Joachim Hornegger¹,
and Dorin Comaniciu²

Presented By:
Anmol Sharma³

¹Friedrich Alexander-Universitat Erlangen-Nurnberg

²Siemens Healthineers

³Medical Image Analysis Laboratory
School of Computing Science
Simon Fraser University
Burnaby, Canada

March 14, 2018

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Problem Statement

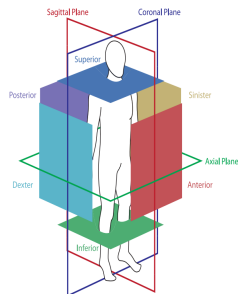
Data and Objects of Interest

- ▶ Detection of anatomical structures in CT scans.

Problem Statement

Data and Objects of Interest

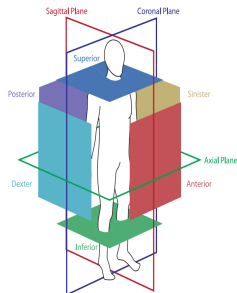
- Detection of anatomical structures in CT scans.



Problem Statement

Data and Objects of Interest

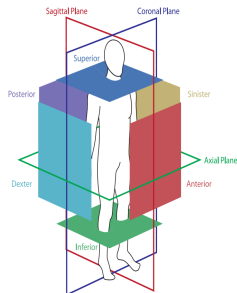
- ▶ Detection of anatomical structures in CT scans.
- ▶ Anatomical structures include:



Problem Statement

Data and Objects of Interest

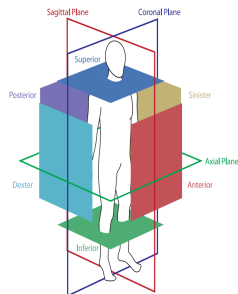
- ▶ Detection of anatomical structures in CT scans.
- ▶ Anatomical structures include:
 - ▶ Left and right kidneys.



Problem Statement

Data and Objects of Interest

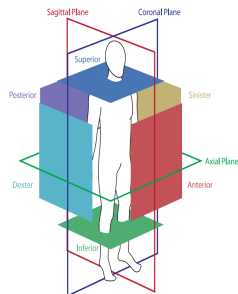
- ▶ Detection of anatomical structures in CT scans.
- ▶ Anatomical structures include:
 - ▶ Left and right kidneys.
 - ▶ Left and right hip bones.



Problem Statement

Data and Objects of Interest

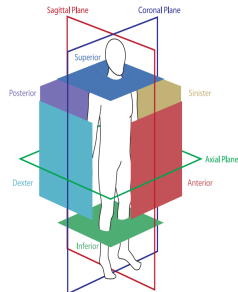
- ▶ Detection of anatomical structures in CT scans.
- ▶ Anatomical structures include:
 - ▶ Left and right kidneys.
 - ▶ Left and right hip bones.
 - ▶ Bronchial bifurcation.



Problem Statement

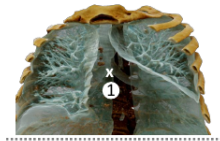
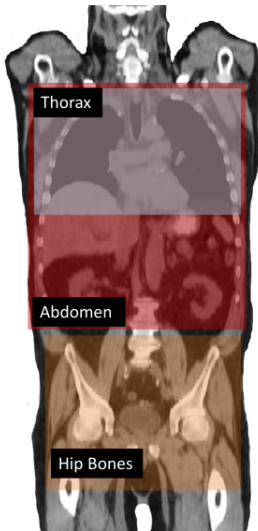
Data and Objects of Interest

- ▶ Detection of anatomical structures in CT scans.
- ▶ Anatomical structures include:
 - ▶ Left and right kidneys.
 - ▶ Left and right hip bones.
 - ▶ Bronchial bifurcation.
 - ▶ And 5 more anatomical structures...

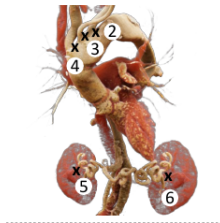


Problem Statement

Data and Objects of Interest



Bronchial
Bifurcation



Aortic Arch
Bifurcations



Kidneys

Hip Bones

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

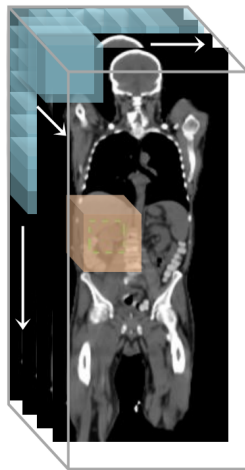
- Learning Multi-Scale Search Strategies

Results

Appendix

Previous Work

Scanning Based Systems



Exhaustive Search

Figure: Systems that scan the image exhaustively or through a smarter region proposal based method.

Previous Work

End-to-End Systems

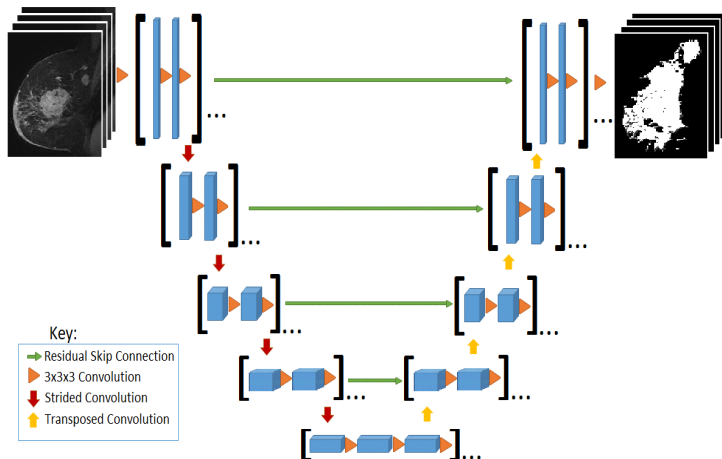


Figure: Systems which output the same sized localization or segmentation maps.

Previous Work

Regression-based Systems

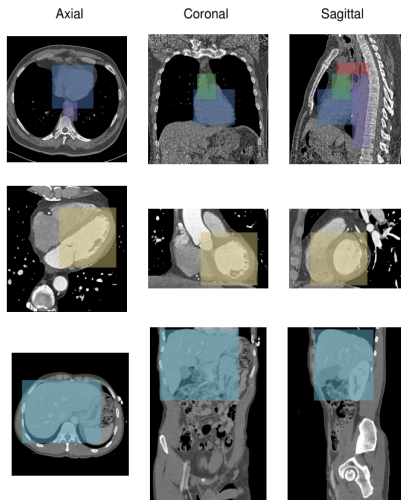


Figure: Systems directly output the position of the structure of interest.

Previous Work

Atlas-based Systems

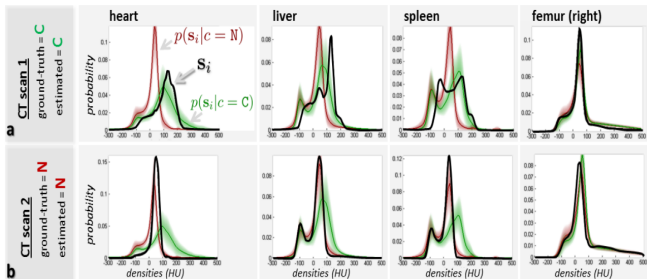


Figure: Systems that search the image using a learnt atlas. Also called Generative Models.

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a

.

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a **behaviour learning task for an artificial agent**.

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a **behaviour learning task for an artificial agent**.
- ▶ Model the

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a **behaviour learning task for an artificial agent**.
- ▶ Model the **anatomy of appearance**

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a **behaviour learning task for an artificial agent**.
- ▶ Model the **anatomy of appearance + object search**

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a **behaviour learning task for an artificial agent**.
- ▶ Model the **anatomy of appearance + object search** in a **unified framework**.

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a **behaviour learning task for an artificial agent**.
- ▶ Model the **anatomy of appearance + object search** in a **unified framework**.
 - ▶ Combine the knowledge of how the object looks

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a **behaviour learning task for an artificial agent**.
- ▶ Model the **anatomy of appearance + object search in a unified framework**.
 - ▶ Combine the knowledge of how the object looks
 - ▶ and where to find it.

Proposed Reformulation Overview

Description

- ▶ Reformulate the structure of interest detection (object detection) problem as a **behaviour learning task for an artificial agent**.
- ▶ Model the **anatomy of appearance + object search** in a **unified framework**.
 - ▶ Combine the knowledge of how the object looks
 - ▶ and where to find it.
- ▶ Model the problem as a **Reinforcement Learning Agent**.

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Reinforcement Learning

In a Nutshell

- ▶ Theory and algorithms uniting multiple fields of machine learning, optimal control theory, psychology and neuroscience.

Reinforcement Learning

In a Nutshell

- ▶ Theory and algorithms uniting multiple fields of machine learning, optimal control theory, psychology and neuroscience.
- ▶ A class of machine learning algorithms that learn by trial-and-error.

Reinforcement Learning

In a Nutshell

- ▶ Theory and algorithms uniting multiple fields of machine learning, optimal control theory, psychology and neuroscience.
- ▶ A class of machine learning algorithms that learn by trial-and-error.
- ▶ Success is measured by a scalar **reward** signal.

Reinforcement Learning

In a Nutshell

- ▶ Theory and algorithms uniting multiple fields of machine learning, optimal control theory, psychology and neuroscience.
- ▶ A class of machine learning algorithms that learn by trial-and-error.
- ▶ Success is measured by a scalar **reward** signal.
- ▶ Goal: select **actions** to **maximize future reward**.

Reinforcement Learning

Atari Example

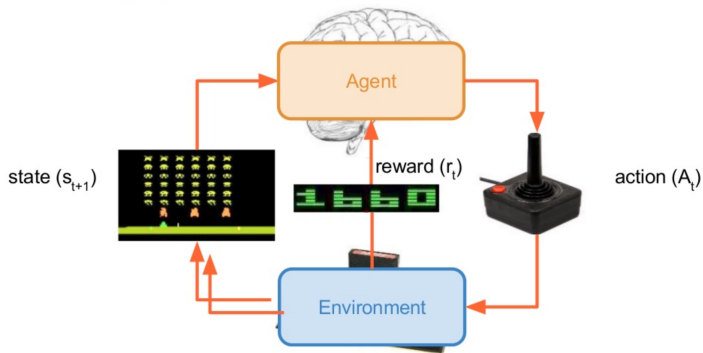


Figure: Agent-Environment interaction

Reinforcement Learning

Game Playing Agent Example



Figure: Mario trying to figure out life (like most grad students)

Reinforcement Learning

Game Playing Agent Example



Figure: Mario trying to figure out life (like most grad students)

- Think of the agent as a robot playing mario.

Reinforcement Learning

Game Playing Agent Example



Figure: Mario trying to figure out life (like most grad students)

- ▶ Think of the agent as a robot playing mario.
- ▶ Should the robot press jump button? (To take the big star?)

Outline

Problem Statement

Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

Reinforcement Learning

Building Blocks of an RL Agent

Method

Deep Reinforcement Learning

Anatomy Detection: Behaviour Learning Problem

Modeling the Agent

Drawing Parallels

Multi-Scale Search Strategy

Limitation

Learning Multi-Scale Search Strategies

Results

Appendix

Policy Function

Predict What Action to Take

$$\pi \left(\text{Current State} \right) = \text{"Jump"}$$


Policy Function

Predict What Action to Take

$$\pi \left(\text{Current State} \right) = \text{"Jump"}$$


- The state is input to the policy function as an image.

Policy Function

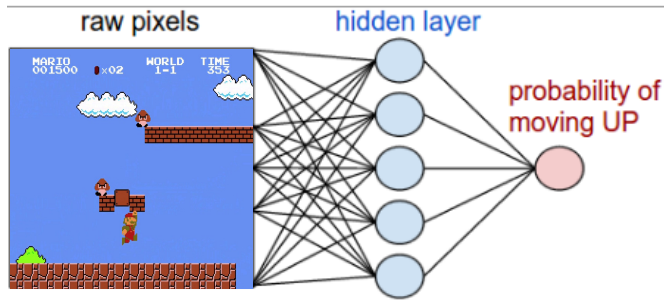
Predict What Action to Take

$$\pi \left(\text{Current State} \right) = \text{"Jump"}$$


- ▶ The state is input to the policy function as an image.
- ▶ That is, a “screenshot” of the game at a time stamp t is called the “state” at that time.

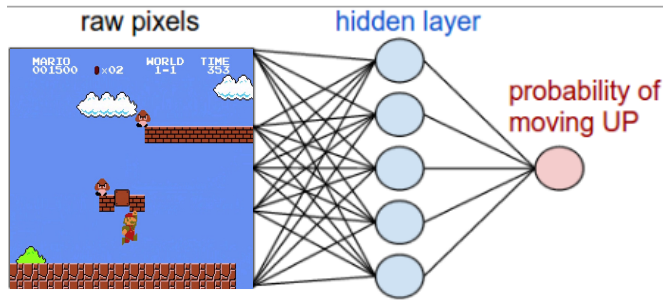
Policy Function

Representation



Policy Function

Representation



- Policy is represented as a **Neural Network** which takes the state (in pixels) as input and outputs probability of an action.

Reward Function

Compute Immediate Reward for an Action

$$\mathcal{R}\left(\begin{array}{c} \text{Current State} \\ \text{Current Action} \end{array}, \text{"Jump"}\right) = 100$$

Reward Function

Compute Immediate Reward for an Action

$$\mathcal{R}\left(\begin{array}{c} \text{Current State} \\ \text{Current Action} \end{array}\right) = 100$$

The equation shows the reward function $\mathcal{R}$ taking a "Current State" and a "Current Action" as inputs. The "Current State" is represented by a screenshot of a Super Mario Bros. game level. The "Current Action" is labeled as "Jump". The result of the function is 100.

- Reward function calculates **immediate reward**.

Reward Function

Compute Immediate Reward for an Action

$$\mathcal{R}\left(\begin{array}{c} \text{Current State} \\ \text{Current Action} \end{array}\right) = 100$$

The image shows a Super Mario Bros. game screen. The top status bar displays "MARIO 001500", a coin count of "1-02", "WORLD 1-1", and "TIME 3:55". The game scene features a brick floor, a green hill, a blue sky with clouds, and a Goomba enemy. Mario is in the air, having just jumped over a brick block.

- ▶ Reward function calculates **immediate reward**.
- ▶ For Atari games, it is predefined in the emulator (games decide how much reward you get).

Reward Function

Compute Immediate Reward for an Action

$$\mathcal{R}\left(\begin{array}{c} \text{Current State} \\ \text{Current Action} \end{array}\right) = 100$$

The image shows a Super Mario Bros. game screen. The top status bar displays 'MARIO 001500', a coin count of '02', 'WORLD 1-1', and 'TIME 3:15'. The game scene features a brick floor, a green hill, a blue sky with clouds, and a brick platform. Mario is on the ground, and a Goomba enemy is on the platform. The action 'Jump' is indicated by the text 'Jump' in quotes.

- ▶ Reward function calculates **immediate reward**.
- ▶ For Atari games, it is predefined in the emulator (games decide how much reward you get).
- ▶ For real world custom problems, you have to define one yourself.

Value Function

Predict Future Reward for an Action-State

$$V \left(\begin{array}{c} \text{Current State} \\ \text{Current Action} \end{array} \right) \approx 5000$$

The image shows a Super Mario Bros. game screen as the "Current State". The screen displays Mario at the bottom, a Goomba enemy on a brick platform, and a Piranha Plant on a higher platform. The top of the screen shows the score (001500), a coin count (02), the world (1-1), and the time (353).

Value Function

Predict Future Reward for an Action-State

$$V \left(\begin{array}{c} \text{Current State} \\ \text{Current Action} \end{array} \right) \approx 5000$$

The image shows a Super Mario Bros. game screen as the "Current State". The screen displays Mario on a brick platform, with a Goomba enemy nearby. The top of the screen shows the game status: "MARIO 001500", "WORLD 1-1", and "TIME 353". The "Current Action" is labeled as "Jump".

- Value function predicts future reward.

Trajectory

Path Taken by Agent in Environment

Trajectory

Path Taken by Agent in Environment

$$\tau = \left(\begin{array}{c} \text{[Screenshot 1]} \\ \text{"Jump", 10} \end{array}, \begin{array}{c} \text{[Screenshot 2]} \\ \text{"Forward", 10} \end{array}, \begin{array}{c} \text{[Screenshot 3]} \\ \text{...} \end{array}, \begin{array}{c} \text{[Screenshot 4]} \\ \text{"Jump", 10} \end{array} \right)$$

$$\tau = (s_0, a_0, r_0, s_1, a_1, r_1, \dots, s_n, a_n, r_n)$$

Total Reward

Total Reward for a Trajectory

- ▶ We need to find the total reward that we received after this whole episode.

Total Reward

Total Reward for a Trajectory

- ▶ We need to find the total reward that we received after this whole episode.
- ▶ It is given as:

Total Reward

Total Reward for a Trajectory

- ▶ We need to find the total reward that we received after this whole episode.
- ▶ It is given as:

$$R(\tau) = \sum_{i=0}^{H-1} \gamma^i * r_i$$

Total Reward

Total Reward for a Trajectory

- ▶ We need to find the total reward that we received after this whole episode.
- ▶ It is given as:

$$R(\tau) = \sum_{i=0}^{H-1} \gamma^i * r_i$$

- ▶ Where γ is the “discount factor”.

Total Reward

Total Reward for a Trajectory

- ▶ We need to find the total reward that we received after this whole episode.
- ▶ It is given as:

$$R(\tau) = \sum_{i=0}^{H-1} \gamma^i * r_i$$

- ▶ Where γ is the “discount factor”.
 - ▶ Larger the value of γ , greater the influence of later rewards on the return.

Total Reward

Total Reward for a Trajectory

$$\mathcal{T} = \left(\begin{array}{c} \text{Gameplay Screenshot} \\ \text{"Jump", 10} \end{array}, \begin{array}{c} \text{Gameplay Screenshot} \\ \text{"Forward", 10} \end{array}, \begin{array}{c} \text{Gameplay Screenshot} \\ \text{, , , ,} \end{array}, \begin{array}{c} \text{Gameplay Screenshot} \\ \text{"Jump", 10} \end{array} \right)$$

$\mathcal{R} = \gamma^0 * 10 + \gamma^1 * 10 + \dots + \gamma^n * 10$

The diagram illustrates the total reward for a trajectory. A sequence of game states (screenshots) is shown, each associated with an action and a reward. The first state shows a character jumping, resulting in a reward of 10. The second state shows a character moving forward, also resulting in a reward of 10. The sequence continues with several more states, each resulting in a reward of 10. The total reward is calculated as the sum of discounted rewards: $\mathcal{R} = \gamma^0 * 10 + \gamma^1 * 10 + \dots + \gamma^n * 10$.

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning**

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Method

Deep Learning + Reinforcement Learning

Method

Deep Learning + Reinforcement Learning

- ▶ Proposed to apply **Deep Reinforcement Learning** for the task of detecting anatomical structures in 3D CT Scans.

Method

Deep Learning + Reinforcement Learning

- ▶ Proposed to apply **Deep Reinforcement Learning** for the task of detecting anatomical structures in 3D CT Scans.
- ▶ Combination of:
 - ▶ **Convolutional Neural Network** as the **Policy Function** and,

Method

Deep Learning + Reinforcement Learning

- ▶ Proposed to apply **Deep Reinforcement Learning** for the task of detecting anatomical structures in 3D CT Scans.
- ▶ Combination of:
 - ▶ **Convolutional Neural Network** as the **Policy Function** and,
 - ▶ Cognitive Model of Reinforcement Learning Framework.

Method

Deep Learning + Reinforcement Learning

- ▶ Proposed to apply **Deep Reinforcement Learning** for the task of detecting anatomical structures in 3D CT Scans.
- ▶ Combination of:
 - ▶ **Convolutional Neural Network** as the **Policy Function** and,
 - ▶ Cognitive Model of Reinforcement Learning Framework.
- ▶ A system capable of both:

Method

Deep Learning + Reinforcement Learning

- ▶ Proposed to apply **Deep Reinforcement Learning** for the task of detecting anatomical structures in 3D CT Scans.
- ▶ Combination of:
 - ▶ **Convolutional Neural Network** as the **Policy Function** and,
 - ▶ Cognitive Model of Reinforcement Learning Framework.
- ▶ A system capable of both:
 - ▶ Identifying an anatomical structure

Method

Deep Learning + Reinforcement Learning

- ▶ Proposed to apply **Deep Reinforcement Learning** for the task of detecting anatomical structures in 3D CT Scans.
- ▶ Combination of:
 - ▶ **Convolutional Neural Network** as the **Policy Function** and,
 - ▶ Cognitive Model of Reinforcement Learning Framework.
- ▶ A system capable of both:
 - ▶ Identifying an anatomical structure
 - ▶ Able to find the structure in a full 3D scan.

Method

Deep Learning + Reinforcement Learning

- ▶ Proposed to apply **Deep Reinforcement Learning** for the task of detecting anatomical structures in 3D CT Scans.
- ▶ Combination of:
 - ▶ **Convolutional Neural Network** as the **Policy Function** and,
 - ▶ Cognitive Model of Reinforcement Learning Framework.
- ▶ A system capable of both:
 - ▶ Identifying an anatomical structure
 - ▶ Able to find the structure in a full 3D scan.

Outline

Problem Statement

Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

Reinforcement Learning

Building Blocks of an RL Agent

Method

Deep Reinforcement Learning

Anatomy Detection: Behaviour Learning Problem

Modeling the Agent

Drawing Parallels

Multi-Scale Search Strategy

Limitation

Learning Multi-Scale Search Strategies

Results

Appendix

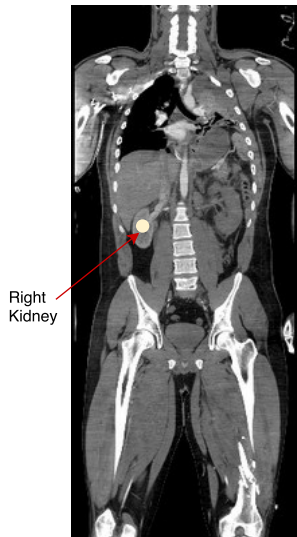
Problem Definition

Visualization



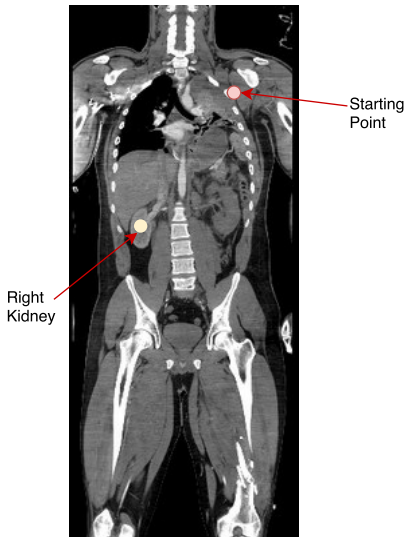
Problem Definition

Visualization



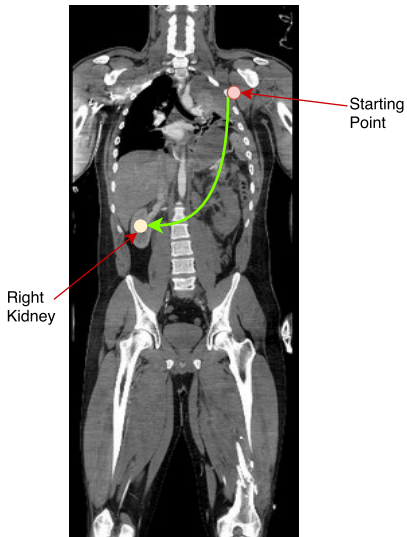
Problem Definition

Visualization



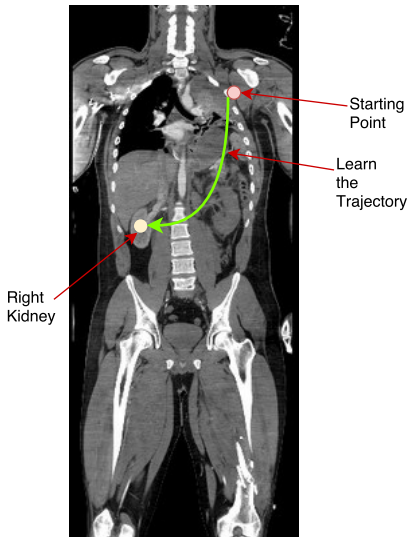
Problem Definition

Visualization



Problem Definition

Visualization



Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent**

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Modeling the Agent

Agent as a controller of Moving Cube



Learned Search-Path

Figure: The search is performed by a moving cube with $\vec{p}_k$ as its centre.

Modeling the Agent

Agent as a controller of Moving Cube



Learned Search-Path

Figure: The search is performed by a moving cube with $\vec{p}_k$ as its centre.

Mario == Cube

Modeling the Agent

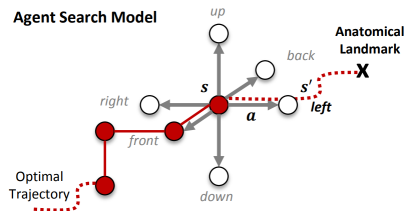
Discrete Action Space

- ▶ The agent is allowed to take a total of 6 actions.

Modeling the Agent

Discrete Action Space

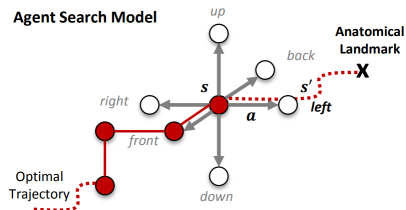
- ▶ The agent is allowed to take a total of 6 actions.



Modeling the Agent

Discrete Action Space

- ▶ The agent is allowed to take a total of 6 actions.
- ▶ It can move up, down, left, right, front or back of the 3D volume.



Reward Function

Defining Reward Function

- ▶ The reward function for this problem is simply:

Reward Function

Defining Reward Function

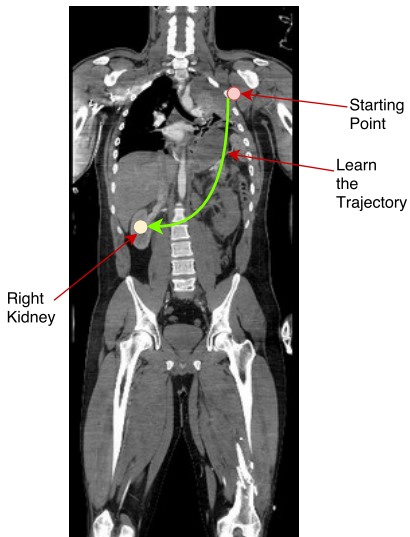
- ▶ The reward function for this problem is simply:

$$r_k = ||\vec{p}_k - \vec{p}_{GT}||$$

- ▶ Euclidean Distance between the current position, and the ground truth position.

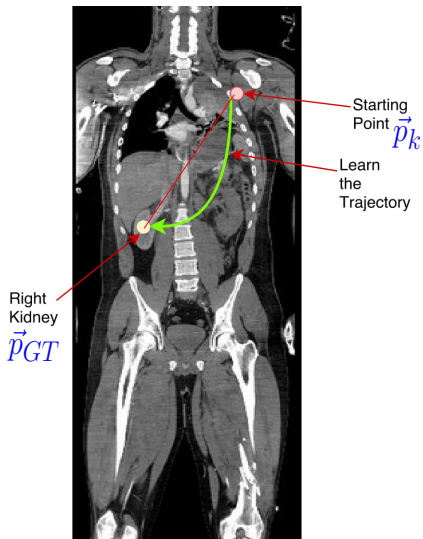
Reward Function

Defining Reward Function



Reward Function

Defining Reward Function



Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels**

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Drawing Parallels

Game Playing RL vs Object Detection RL: Policy

Connect the earlier terminology to this proposed terminology.

Drawing Parallels

Game Playing RL vs Object Detection RL: Policy

Connect the earlier terminology to this proposed terminology.

$$\pi \left(\begin{array}{c} \text{Atari Game Frame} \\ \text{Current State} \end{array} \right) = \text{"Jump"} \quad \left| \quad \pi \left(\begin{array}{c} \text{3D Medical Scan Frame} \\ \text{Current State} \end{array} \right) = \text{"Up"}$$

In both cases, Policy Function is a Convolutional Neural Network (2D for Atari and 3D for CT).

Drawing Parallels

Game Playing RL vs Object Detection RL: Reward

Connect the earlier terminology to this proposed terminology.

Drawing Parallels

Game Playing RL vs Object Detection RL: Reward

Connect the earlier terminology to this proposed terminology.

$$\begin{array}{ccc} \mathcal{R}\left(\underbrace{\text{Current State}}_{\text{Current State}}, \text{"Jump"}\right) = 10 & \text{---} & \mathcal{R}\left(\underbrace{\text{Current State}}_{\text{Current State}}, \underbrace{\text{New State}}_{\text{New State}}, \underbrace{\text{"Left"}}_{\text{Action}}\right) = \|\vec{p}_{new} - \vec{p}_{GT}\| \\ \text{Reward function defined by Environment} & & \text{Reward function defined by User} \end{array}$$

The reward function in this work was defined as distance between the agent's position and ground truth position.

Drawing Parallels

Game Playing RL vs Object Detection RL: Trajectory

Connect the earlier terminology to this proposed terminology.

Drawing Parallels

Game Playing RL vs Object Detection RL: Trajectory

Connect the earlier terminology to this proposed terminology.

$$\mathcal{T} = \left(\begin{array}{c} \text{Game Playing RL} \\ \text{Object Detection RL} \end{array} \right)$$

Game Playing RL Trajectory:

$$\mathcal{T} = \left(\begin{array}{c} \text{Image 1} \\ \text{Image 2} \\ \text{Image 3} \\ \vdots \\ \text{Image N} \end{array} \right)$$

Object Detection RL Trajectory:

$$\mathcal{T} = \left(\begin{array}{c} \text{Image 1} \\ \text{Image 2} \\ \text{Image 3} \\ \vdots \\ \text{Image N} \end{array} \right)$$

The trajectory vector stores all “states” or 3D patch volumes that it visited, and the corresponding actions that the agent took.

Optimization Function

Traditional vs Deep

- ▶ In traditional RL, we use

$$Q^*(s, a) = \mathbb{E}_{s'} (r + \gamma \max_{a'} Q^*(s', a'))$$

- ▶ As objective function, also called the Bellman Equation.
- ▶ However, due to high dimensional nature of state space (state if defined in terms of position in the 3D image), this function is unsuitable.

Optimization Function

Traditional vs Deep

- ▶ In traditional RL, we use

$$Q^*(s, a) = \mathbb{E}_{s'} (r + \gamma \max_{a'} Q^*(s', a'))$$

- ▶ As objective function, also called the Bellman Equation.
- ▶ However, due to high dimensional nature of state space (state if defined in terms of position in the 3D image), this function is unsuitable.
- ▶ **Solution?**

Optimization Function

Traditional vs Deep

- ▶ In traditional RL, we use

$$Q^*(s, a) = \mathbb{E}_{s'} (r + \gamma \max_{a'} Q^*(s', a'))$$

- ▶ As objective function, also called the Bellman Equation.
- ▶ However, due to high dimensional nature of state space (state if defined in terms of position in the 3D image), this function is unsuitable.
- ▶ **Solution?**
- ▶ Use a Deep Convolutional Neural Network as a non-linear parameterization of Q^* .

Optimization Function

Traditional vs Deep

- ▶ In traditional RL, we use

$$Q^*(s, a) = \mathbb{E}_{s'} (r + \gamma \max_{a'} Q^*(s', a'))$$

- ▶ As objective function, also called the Bellman Equation.
- ▶ However, due to high dimensional nature of state space (state if defined in terms of position in the 3D image), this function is unsuitable.
- ▶ **Solution?**
- ▶ Use a Deep Convolutional Neural Network as a non-linear parameterization of Q^* .
- ▶ This is called Deep Q-Network (DQN).

Optimization Function

Bellman Optimality with MSE

- ▶ A Deep Q-Network can be trained in a RL setup using iterative approach.

Optimization Function

Bellman Optimality with MSE

- ▶ A Deep Q-Network can be trained in a RL setup using iterative approach.
- ▶ Minimizes the mean squared error based on the Bellman optimality equation.

Optimization Function

Bellman Optimality with MSE

- ▶ A Deep Q-Network can be trained in a RL setup using iterative approach.
- ▶ Minimizes the mean squared error based on the Bellman optimality equation.

$$\hat{\theta}^{(i)} = \underset{\theta^{(i)}}{\operatorname{argmin}} \mathbb{E}_{s,a,r,s'} [(y - Q(s, a; \theta^{(i)}))^2]$$

Optimization Function

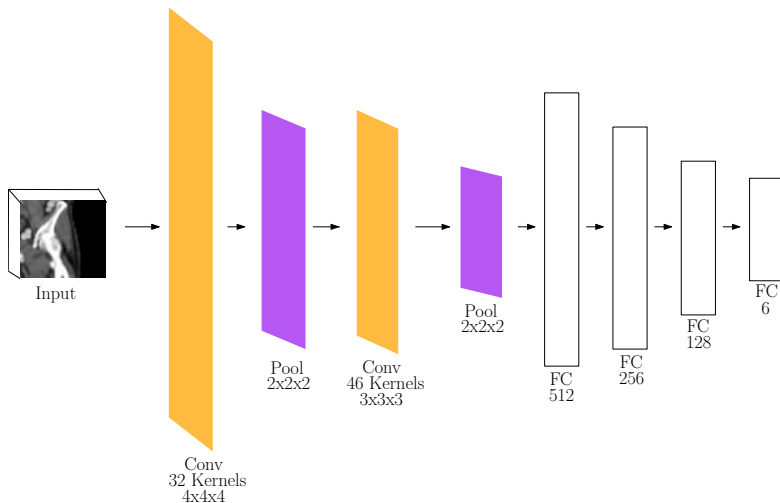
Bellman Optimality with MSE

- ▶ A Deep Q-Network can be trained in a RL setup using iterative approach.
- ▶ Minimizes the mean squared error based on the Bellman optimality equation.

$$\hat{\theta}^{(i)} = \underset{\theta^{(i)}}{\operatorname{argmin}} \mathbb{E}_{s,a,r,s'} [(y - Q(s, a; \theta^{(i)}))^2]$$

- ▶ where $y = r + \gamma \max_{a'} Q(s', a'; \theta^{(i)})$

Network Architecture



Network Architecture

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Outline

Problem Statement

Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

Reinforcement Learning

Building Blocks of an RL Agent

Method

Deep Reinforcement Learning

Anatomy Detection: Behaviour Learning Problem

Modeling the Agent

Drawing Parallels

Multi-Scale Search Strategy

Limitation

Learning Multi-Scale Search Strategies

Results

Appendix

Major Limitation

Major Limitation

- ▶ An obvious **limitation** of the approach is the fixed box size of the agent, as shown:

Major Limitation

- ▶ An obvious **limitation** of the approach is the fixed box size of the agent, as shown:
- ▶ What if the size of object is much smaller than the box?



Learned Search-Path

Major Limitation

- ▶ An obvious **limitation** of the approach is the fixed box size of the agent, as shown:
- ▶ What if the size of object is much smaller than the box?
 - ▶ Localization is useless if the box is too large.



Learned Search-Path

Major Limitation

- ▶ An obvious **limitation** of the approach is the fixed box size of the agent, as shown:
- ▶ What if the size of object is much smaller than the box?
 - ▶ Localization is useless if the box is too large.
- ▶ Need to make the system **scale-invariant**.



Learned Search-Path

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

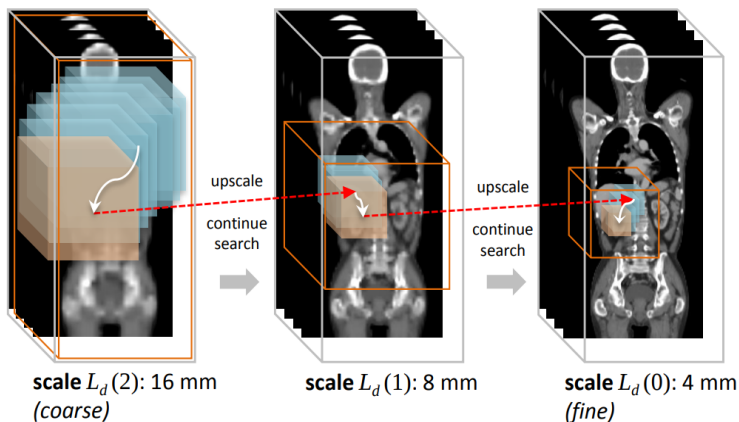
- Learning Multi-Scale Search Strategies

Results

Appendix

Multi-Scale Search Strategies

Visualization



Multi-Scale Search Strategies

Solution to Scale-invariance

- ▶ **Solution:** Search on multiple scales of the input 3D volume.
- ▶ Obtain a finite number of scaled 3D volumes, starting from coarsest (16mm) to finest (4mm).
- ▶ Search algorithm:
 - ▶ Start the search from the coarsest scale (lowest resolution).
 - ▶ If the agent **converges**¹, change to finer (higher resolution) scale, and start the search again from the convergence point.
 - ▶ Repeat until convergence on finest scale is reached.

¹discussed in the next slide

Multi-Scale Search Strategies

Convergence Criteria

- ▶ The convergence criteria for the agent answers the question:
When does the agent know it has found the object of interest?
- ▶ Answer: Detect small oscillatory-like cycles in the agent's behaviour.
- ▶ If a cycle is detected, save the position, and move to next scale level to continue search.
 - ▶ If the scale is finest, return the current position as the position of object of interest.

Multi-Scale Search Strategies

Increased Complexity in Scale Space

- ▶ **However**
- ▶ Searching through multiple scale spaces increases the complexity in learning of policy function (CNN).
- ▶ The CNN would have to learn
 - ▶ how to distinguish an object of interest from background
 - ▶ how to navigate the image space by taking actions
 - ▶ how to do the above 2 tasks in multiple scales.
- ▶ This would make the learning intractable.
- ▶ **Solution:** Authors propose to use multiple CNNs for multiple scales.

Multi-Scale Search Strategies

Increased Complexity in Scale Space



Separate CNNs for separate scale 3D images.

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Comparison

Algorithms compared against

Comparison

Algorithms compared against

- ▶ The proposed method was compared against:

Comparison

Algorithms compared against

- ▶ The proposed method was compared against:
 - ▶ Probabilistic Boosting Trees

Comparison

Algorithms compared against

- ▶ The proposed method was compared against:
 - ▶ Probabilistic Boosting Trees
 - ▶ Extremely Randomized Trees with Hough Regression

Comparison

Algorithms compared against

- ▶ The proposed method was compared against:
 - ▶ Probabilistic Boosting Trees
 - ▶ Extremely Randomized Trees with Hough Regression
 - ▶ Overfeat, adapted to 3D data

Comparison

Algorithms compared against

- ▶ The proposed method was compared against:
 - ▶ Probabilistic Boosting Trees
 - ▶ Extremely Randomized Trees with Hough Regression
 - ▶ Overfeat, adapted to 3D data
 - ▶ 3D Deep Learning with Filter Decomposition

Comparison

Algorithms compared against

- ▶ The proposed method was compared against:
 - ▶ Probabilistic Boosting Trees
 - ▶ Extremely Randomized Trees with Hough Regression
 - ▶ Overfeat, adapted to 3D data
 - ▶ 3D Deep Learning with Filter Decomposition
 - ▶ Scanning with Cascaded Sparse-Adaptive DNN

► **Failure Percentage Rate**

- If the agent doesn't come within:
 - 30mm for the center of each kidney.
 - 10mm for the remaining landmarks
- It is deemed a failure of detection.

Metrics

- ▶ **Failure Percentage Rate**

- ▶ If the agent doesn't come within:
 - ▶ 30mm for the center of each kidney.
 - ▶ 10mm for the remaining landmarks
- ▶ It is deemed a failure of detection.

- ▶ **Accuracy**

- ▶ According to the number of times the agent detected the structure successfully (using above definition).

Metrics

- ▶ **Failure Percentage Rate**

- ▶ If the agent doesn't come within:
 - ▶ 30mm for the center of each kidney.
 - ▶ 10mm for the remaining landmarks
- ▶ It is deemed a failure of detection.

- ▶ **Accuracy**

- ▶ According to the number of times the agent detected the structure successfully (using above definition).

- ▶ **Speed**

- ▶ Raw speed of detection on their hardware platform

Quantitative Results

Performance Compared to Other Published Works

Solution	Dataset Size (Data/Patients)	Accuracy (mm)	Speed (seconds)
Zhan <i>et al.</i> [27]	18/18 CT	4.72	4
Fenchel <i>et al.</i> [35]	31/31 MR	22.4	20
Criminisi <i>et al.</i> [12]	100/– CT	17.60	1
Pauly <i>et al.</i> [32]	33/33 MR	14.95	0.8
Cuingnet <i>et al.</i> [11]	233/89 CT	10.5	2.8
Donner <i>et al.</i> [10]	20/20 CT	5.25	120
Criminisi <i>et al.</i> [31]	400/– CT	13.50	4
Chu <i>et al.</i> [30]	10/10 CT	1.90 ¹	30
Potesil <i>et al.</i> [37]	83/83 CT	4.70	N/A
de Vos <i>et al.</i> [24]	100/– CT	4.80	10
Ours	1487/532 CT	4.19²	0.061

¹ Evaluated only on vertebrae localization with strong priors.

² With no failures of clinical significance. All other solutions did not provide any information in this respect.

Quantitative Results

Speed Compared to Methods Discussed in the Previous Slide

DETECTION TIME per Volume [seconds]

Method	Platform	Median	Min	Max
PBT [3]	CPU 8-core	1.051	0.14	16.20
ExtRTrees [10]	CPU 8-core	4.714	0.272	44.535
Overfeat ¹ [22]	GPU Titan X	2.175	0.331	14.86
3D-DL [18]	CPU 10-core	57.034	7.161	548.23
SADNN ² [1]	GPU GTX 1080	0.471	0.064	3.029
Ours	CPU 8-core	0.061	0.035	0.155
Ours	GPU Titan X	0.033	0.018	0.085

^{1,2} Evaluated on finest isotropic resolution of 4 mm.

Quantitative Results

Individual Anatomy Detection Comparison with Other Methods

Structure Type	Landmark	Method	Failed Cases			Accuracy (excl. failed cases)	
			% Failed	Median	Max	Mean $\pm$ STD	Median
Bone	Hip Bone Right Front Corner	PBT [3]	4.35%	38.89	141.88	3.05 ± 1.60	2.77
		ExtRTrees [10]	8.07%	20.27	460.22	5.01 ± 2.40	4.95
		Overfeat [22]	9.31%	35.64	231.29	4.55 ± 1.98	4.31
		3D-DL [18]	0.62%	10.17	10.17	2.84 ± 1.36	2.53
		SADNN [1]	0%	–	–	3.50 ± 1.63	3.37
		Ours	0%	–	–	2.80 ± 1.46	2.53
	Hip Bone Left Front Corner	PBT [3]	3.75%	14.30	28.57	4.03 ± 1.79	3.86
		ExtRTrees [10]	6.25%	13.43	17.44	5.13 ± 2.78	5.08
		Overfeat [22]	3.75%	127.52	242.05	4.19 ± 1.80	4.15
		3D-DL [18]	2.50%	262.81	513.81	3.14 ± 1.49	2.98
		SADNN [1]	1.25%	12.57	12.57	4.44 ± 1.75	4.43
		Ours	0%	–	–	3.07 ± 2.14	2.85

Critique

Some Critical Observations

- ▶ Authors reimplemented a method without a main subsystem.
- ▶ Runtime comparison was performed on heterogeneous systems.
- ▶ Self imposed accuracy threshold.
 - ▶ If the agent is within 10mm of the body part, it is deemed as successful detection.
 - ▶ No analysis of reducing/increasing this threshold to report varying results.
 - ▶ Although the authors cite clinical reasons to choose this threshold.
- ▶ Starting point of the agent is heuristically chosen to be close to the actual expected location.
 - ▶ Makes sense because searching for right kidney does not have to start from the head.
 - ▶ However, no results reported about runtime speed and accuracy if the chosen point is randomized.

Critique

Some Critical Observations

Pros

- ▶ A brilliantly written paper.
- ▶ Highly polished analysis, clear and concise handling of various apparent limitations.

Conclusion

- ▶ Authors proposed a new formulation of the object detection problem in 3D CT Scans.
- ▶ Posed the problem as Deep Reinforcement Learning based agent learning to navigate the image space and identifying markers.
- ▶ Used Deep Q-Learning with 3D Convolutional Neural Network as policy function.
- ▶ Reported 0% Failure Percentage Rate for their approach.
- ▶ Significant orders of speed-up in detection process, compared to state-of-art methods.

Thank You for listening!

Outline

Problem Statement

- Data and Objects of Interest

Conventional Approaches

Proposed Reformulation of Problem

Background

- Reinforcement Learning

- Building Blocks of an RL Agent

Method

- Deep Reinforcement Learning

- Anatomy Detection: Behaviour Learning Problem

- Modeling the Agent

- Drawing Parallels

Multi-Scale Search Strategy

- Limitation

- Learning Multi-Scale Search Strategies

Results

Appendix

Starting Point of the Agent

How to set agent's starting point?

- ▶ Also, the starting point of the agent is set to a position $\vec{p}_0$ using the expected position of the object of interest.
- ▶ This position is sampled from the training data.
- ▶ For example, for right kidney, it is known that it is on the right side of the abdomen. Hence set the starting position close to that region.

Object Not in Scan?

How to find out if object is not in scan?

Object Not in Scan?

How to find out if object is not in scan?

- ▶ In any object detection framework, the system or agent should be able to handle the case where the object is not present in the image (or scene).

Object Not in Scan?

How to find out if object is not in scan?

- ▶ In any object detection framework, the system or agent should be able to handle the case where the object is not present in the image (or scene).
- ▶ In this work, a very simple heuristic approach is taken to find out if the object is not present in the scene.

Object Not in Scan?

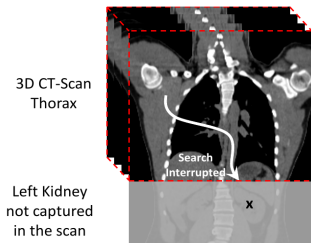
How to find out if object is not in scan?

- ▶ In any object detection framework, the system or agent should be able to handle the case where the object is not present in the image (or scene).
- ▶ In this work, a very simple heuristic approach is taken to find out if the object is not present in the scene.
- ▶ **Approach:** Detect if the agent leaves the image space while scanning.

Object Not in Scan?

How to find out if object is not in scan?

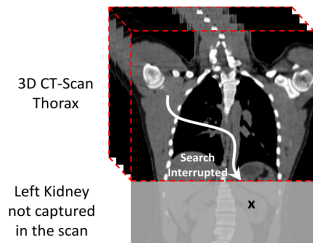
- ▶ In any object detection framework, the system or agent should be able to handle the case where the object is not present in the image (or scene).
- ▶ In this work, a very simple heuristic approach is taken to find out if the object is not present in the scene.
- ▶ **Approach:** Detect if the agent leaves the image space while scanning.



Object Not in Scan?

How to find out if object is not in scan?

- ▶ In any object detection framework, the system or agent should be able to handle the case where the object is not present in the image (or scene).
- ▶ In this work, a very simple heuristic approach is taken to find out if the object is not present in the scene.
- ▶ **Approach:** Detect if the agent leaves the image space while scanning.
- ▶ Although simple, the approach is shown to work atleast for 90%+ negative cases, which show the robustness.



Exploration vs Exploitation

Overview

- ▶ Learning the action-value function Q^* enables the agent to effectively search for objects in the image.

Exploration vs Exploitation

Overview

- ▶ Learning the action-value function Q^* enables the agent to effectively search for objects in the image.
- ▶ However there is a need to balance exploration by the agent against exploitation.

Exploration vs Exploitation

Overview

- ▶ Learning the action-value function Q^* enables the agent to effectively search for objects in the image.
- ▶ However there is a need to balance exploration by the agent against exploitation.
- ▶ The authors propose to use three techniques to ensure adequate exploration of the environment.

ϵ -greedy Approach

ϵ -greedy Approach

- ▶ A variable $\epsilon \in [0, 1]$ controls the randomness in the exploration.

ϵ -greedy Approach

- ▶ A variable $\epsilon \in [0, 1]$ controls the randomness in the exploration.
- ▶ In simple words:

ϵ -greedy Approach

- ▶ A variable $\epsilon \in [0, 1]$ controls the randomness in the exploration.
- ▶ In simple words:
- ▶ During training, sample the actions from

ϵ -greedy Approach

- ▶ A variable $\epsilon \in [0, 1]$ controls the randomness in the exploration.
- ▶ In simple words:
- ▶ During training, sample the actions from
 - ▶ Uniform Distribution, with probability ϵ , or

ϵ -greedy Approach

- ▶ A variable $\epsilon \in [0, 1]$ controls the randomness in the exploration.
- ▶ In simple words:
- ▶ During training, sample the actions from
 - ▶ Uniform Distribution, with probability ϵ , or
 - ▶ Deterministically from the current policy, with probability $1 - \epsilon$.

ϵ -greedy Approach

- ▶ A variable $\epsilon \in [0, 1]$ controls the randomness in the exploration.
- ▶ In simple words:
- ▶ During training, sample the actions from
 - ▶ Uniform Distribution, with probability ϵ , or
 - ▶ Deterministically from the current policy, with probability $1 - \epsilon$.
- ▶ This value is annealed linearly from 1.0 to 0.05 during the training.

Experience Replay

Experience Replay

- ▶ Ensure training stability by decorrelating the training samples using **experience replay**.

Experience Replay

- ▶ Ensure training stability by decorrelating the training samples using **experience replay**.
- ▶ During training, memorize the trajectories in $M = [T_1, T_2, T_3 \dots]$.

Experience Replay

- ▶ Ensure training stability by decorrelating the training samples using **experience replay**.
- ▶ During training, memorize the trajectories in $M = [T_1, T_2, T_3 \dots]$.
- ▶ Uniformly sample from this list of trajectories.

Experience Replay

- ▶ Ensure training stability by decorrelating the training samples using **experience replay**.
- ▶ During training, memorize the trajectories in $M = [T_1, T_2, T_3 \dots]$.
- ▶ Uniformly sample from this list of trajectories.
- ▶ Use those to train the policy (deep convolutional neural network).

Adaptive Episode Length

Adaptive Episode Length

- ▶ Further accelerate the training using an adaptive episode length.

Adaptive Episode Length

- ▶ Further accelerate the training using an adaptive episode length.
- ▶ During training, linearly increase the number of trajectories sampled from the memory, and train using those.

Adaptive Episode Length

- ▶ Further accelerate the training using an adaptive episode length.
- ▶ During training, linearly increase the number of trajectories sampled from the memory, and train using those.
- ▶ That is, don't gather new experience all through the training.

Adaptive Episode Length

- ▶ Further accelerate the training using an adaptive episode length.
- ▶ During training, linearly increase the number of trajectories sampled from the memory, and train using those.
- ▶ That is, don't gather new experience all through the training.
- ▶ As the training progresses, use more experiences of the past rather than making new ones.

Adaptive Episode Length

- ▶ Further accelerate the training using an adaptive episode length.
- ▶ During training, linearly increase the number of trajectories sampled from the memory, and train using those.
- ▶ That is, don't gather new experience all through the training.
- ▶ As the training progresses, use more experiences of the past rather than making new ones.
- ▶ This reduces training time by 30%.