

IMPLEMENTASI REINFORCEMENT LEARNING UNTUK MENENTUKAN TINGKAT KECERDASAN NON-PLAYER- CHARACTER PADA PERMAINAN BALAP MOBIL

Skripsi



Oleh

**YUSTINUS ARJUNA PURNAMA PUTRA
71110164**

PROGRAM STUDI TEKNIK INFORMATIKA FAKULTAS TEKNOLOGI INFORMASI
UNIVERSITAS KRISTEN DUTA WACANA
2016

IMPLEMENTASI REINFORCEMENT LEARNING UNTUK MENENTUKAN TINGKAT KECERDASAN NON-PLAYER- CHARACTER PADA PERMAINAN BALAP MOBIL

Skripsi



Diajukan kepada Program Studi Teknik Informatika Fakultas Teknologi Informasi
Universitas Kristen Duta Wacana
Sebagai Salah Satu Syarat dalam Memperoleh Gelar
Sarjana Komputer

Disusun oleh

**YUSTINUS ARJUNA PURNAMA PUTRA
7111016**

PROGRAM STUDI TEKNIK INFORMATIKA FAKULTAS TEKNOLOGI INFORMASI
UNIVERSITAS KRISTEN DUTA WACANA
2016

PERNYATAAN KEASLIAN SKRIPSI

Saya menyatakan dengan sesungguhnya bahwa skripsi dengan judul:

IMPLEMENTASI REINFORCEMENT LEARNING UNTUK MENENTUKAN TINGKAT KECERDASAN NON-PLAYER- CHARACTER PADA PERMAINAN BALAP MOBIL

yang saya kerjakan untuk melengkapi sebagian persyaratan menjadi Sarjana Komputer pada pendidikan Sarjana Program Studi Teknik Informatika Fakultas Teknologi Informasi Universitas Kristen Duta Wacana, bukan merupakan tiruan atau duplikasi dari skripsi keserjanaan di lingkungan Universitas Kristen Duta Wacana maupun di Perguruan Tinggi atau instansi manapun, kecuali bagian yang sumber informasinya dicantumkan sebagaimana mestinya.

Jika dikemudian hari didapati bahwa hasil skripsi ini adalah hasil plagiasi atau tiruan dari skripsi lain, saya bersedia dikenai sanksi yakni pencabutan gelar keserjanaan saya.

Yogyakarta, 10 Mei 2016



YUSTINUS ARJUNA PURNAMA
PUTRA

HALAMAN PERSETUJUAN

Judul Skripsi : IMPLEMENTASI REINFORCEMENT LEARNING
UNTUK MENENTUKAN TINGKAT
KECERDASAN NON-PLAYER-CHARACTER
PADA PERMAINAN BALAP MOBIL

Nama Mahasiswa : YUSTINUS ARJUNA PURNAMA PUTRA

N I M : 71110164

Matakuliah : Skripsi (Tugas Akhir)

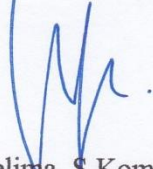
Kode : TIW276

Semester : Genap

Tahun Akademik : 2015/2016

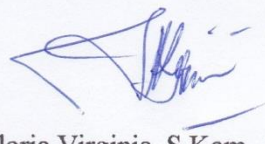
Telah diperiksa dan disetujui di
Yogyakarta,
Pada tanggal 09 Mei 2016

Dosen Pembimbing I



Rosa Delima, S.Kom., M.Kom.

Dosen Pembimbing II



Gloria Virginia, S.Kom., MAI, Ph.D.

HALAMAN PENGESAHAN

IMPLEMENTASI REINFORCEMENT LEARNING UNTUK MENENTUKAN TINGKAT KECERDASAN NON-PLAYER- CHARACTER PADA PERMAINAN BALAP MOBIL

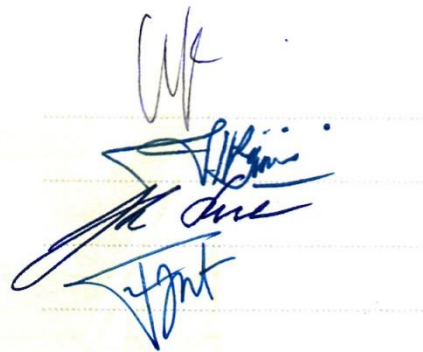
Oleh: YUSTINUS ARJUNA PURNAMA PUTRA / 71110164

Dipertahankan di depan Dewan Penguji Skripsi
Program Studi Teknik Informatika Fakultas Teknologi Informasi
Universitas Kristen Duta Wacana - Yogyakarta
Dan dinyatakan diterima untuk memenuhi salah satu syarat memperoleh gelar
Sarjana Komputer
pada tanggal 27 Mei 2016

Yogyakarta, 24 Juni 2016
Mengesahkan,

Dewan Penguji:

1. Rosa Delima, S.Kom., M.Kom.
2. Gloria Virginia, S.Kom., MAI, Ph.D.
3. Budi Susanto, S.Kom., M.T.
4. Antonius Rachmat C., S.Kom., M.Cs.



Dekan

(Budi Susanto, S.Kom., M.T.)

Ketua Program Studi

(Gloria Virginia, Ph.D.)

UCAPAN TERIMA KASIH

Puji dan syukur penulis panjatkan kepada Tuhan Yang Maha Esa yang telah melimpahkan rahmat dan anugerah sehingga penulis dapat menyelesaikan tugas akhir yang berjudul Implementasi Reinforcement Learning Untuk Menentukan Tingkat Kecerdasan Non-Player-Character Pada Permainan Balap Mobil.

Penulisan laporan ini merupakan kelengkapan dan pemenuhan dari salah satu syarat dalam memperoleh gelar Sarjana Komputer. Selain itu bertujuan melatih mahasiswa untuk dapat menghasilkan suatu karya yang dapat dipertanggungjawabkan secara ilmiah, sehingga dapat bermanfaat bagi penggunaanya.

Dalam menyelesaikan pembuatan program dan laporan Tugas Akhir ini, penulis telah banyak menerima bimbingan, saran, dan masukan dari berbagai pihak, baik secara langsung maupun secara tidak langsung. Untuk itu dengan segala kerendahan hati, pada kesempatan ini penulis menyampaikan ucapan terima-kasih kepada:

1. Rosa Delima, S.Kom., M.Kom. selaku dosen pembimbing I yang telah memberikan bimbingan dengan sabar dan baik kepada penulis.
2. Gloria Virginia, S.Kom., MAI, Ph.D. selaku dosen pembimbing II atas bimbingan, petunjuk, dan saran yang diberikan selama pengerjaan Tugas Akhir ini sejak awal hingga akhir.
3. Dosen-dosen Universitas Kristen Duta Wacana yang telah membantu memberikan saran kepada penulis.
4. Theresia Sunaryati selaku ibu yang telah mendoakan, mendukung sarana penulis dalam melakukan penelitian dan sabar dalam mengingatkan penulis dalam mengerjakan Tugas Akhir.
5. Keluarga Tercinta yang telah mendukung penulis dengan saran dan semangat.

6. Christina Dian Wijayanti yang telah memberikan doa, semangat dan membantu pencarian data penelitian.
7. Yudi Tri Sanjaya atas dukungan dan support dalam mengerjakan Tugas Akhir.
8. Tyfons IT Solution yang memberikan dukungan dan masukkan pada penulis.
9. Semua saudara, teman dan kerabat yang tidak dapat penulis sebutkan satu persatu, sehingga pembuatan Tugas Akhir ini dapat berjalan dengan baik.

Penulis menyadari bahwa program dan laporan Tugas Akhir ini masih jauh dari sempurna. Oleh karena itu, penulis sangat mengharapkan kritik dan saran yang membangun dari pembaca sekalian. Sehingga suatu saat penulis dapat memberikan karya yang lebih baik lagi. Akhir kata penulis ingin meminta maaf bila ada kesalahan baik dalam penyusunan laporan maupun yang pernah penulis lakukan sewaktu membuat program Tugas Akhir. Sekali lagi penulis memohon maaf yang sebesar-besarnya.

Yogyakarta, 8 Mei 2016

Penulis

INTISARI

Implementasi Reinforcement Learning Untuk Menentukan Tingkat Kecerdasan Non-Player-Character Pada Permainan Balap Mobil

Q-learning adalah salah satu jenis pembelajaran mesin yang dilakukan pada lingkungan yang tidak diketahui oleh agen, sehingga agen harus mencoba-coba aksi yang tepat pada suatu *state*. Evaluasi aksi yang dilakukan agen adalah memberikan *reward*. Dalam algoritma ini, agen memperbaharui *Q-table* dengan nilai *reward* tersebut. Tabel tersebut menyimpan nilai action untuk setiap *state* selama pelatihan. Nilai tersebut merepresentasikan bobot kesesuaian aksi pada *state* tertentu.

Dengan mengimplementasikan algoritma *Q-learning*, penulis bermaksud mengembangkan sistem dimana terdapat sejumlah NPC pembalap mobil (selanjutnya dikenal sebagai agen) yang masing-masing memiliki kemampuan untuk mempersepsikan kondisi agen pada lingkungan pembelajaran. Aksi yang dapat dilakukan agen untuk menanggapi *state* yang dipersepsikan adalah berbelok ke kiri, maju lurus dan berbelok ke kanan. Tujuan penelitian ini adalah untuk membuktikan bahwa dengan algoritma *Q-learning*, perbedaan tingkat kecerdasan agen cerdas dapat dibuat dengan cara membedakan variasi lintasan yang dipelajari oleh agen dan membedakan variasi waktu belajar agen. Agen akan terbukti lebih pandai apabila semakin banyak variasi lintasan yang dapat diselesaikan saat pengujian.

Dengan menguji agen cerdas pada lintasan uji dengan variasi berbeda, dapat disimpulkan bahwa dapat ditentukan perbedaan tingkat kecerdasan agen dengan cara membedakan jumlah variasi lintasan yang dipelajari oleh agen cerdas dan cara kedua adalah dengan membedakan lama waktu pembelajaran.

Kata kunci : QLearning, Reinforcement Learning, Mobil Cerdas

DAFTAR ISI

UCAPAN TERIMAKASIH.....	vi
INTISARI.....	viii
BAB I.....	1
PENDAHULUAN	1
1.1 Latar Belakang Masalah.....	1
1.2 Perumusan Masalah.....	2
1.3 Batasan Masalah.....	3
1.4 Tujuan Penelitian.....	3
1.5 Metode Penelitian.....	3
1.6 Sistematika Penulisan.....	4
BAB II.....	6
TINJAUAN PUSTAKA	6
2.1 Tinjauan Pustaka	6
2.2 Landasan Teori	7
2.2.1 <i>Reinforcement Learning</i>	7
2.2.2 <i>Two Line Segment Intersection Test</i>	9
BAB III	11
ANALISIS DAN PERANCANGAN SISTEM	11
3.1 Kebutuhan Sistem.....	11
3.1.1 Kebutuhan Hardware	11
3.1.2 Kebutuhan Software.....	11
3.2 Rancangan Sistem	12
3.2 <i>Flowchart</i> Sistem	13
3.2.1 <i>Flowchart</i> Fungsi Pembuatan Lintasan.....	17
3.2.2 <i>Flowchart</i> Fungsi Pembelajaran Lintasan.....	18
3.2.3 <i>Flowchart</i> Fungsi Pengujian Agen Cerdas	28
3.3 Rancangan Agen Cerdas dan Lingkungan Pembelajaran.....	28

3.4	Rancangan Evaluasi Sistem.....	31
3.5	Rancangan Antar Muka.....	31
BAB IV		35
IMPLEMENTASI DAN ANALISIS SISTEM.....		35
4.1	Implementasi Sistem	35
4.1.1	Konfigurasi Awal.....	35
4.1.2	Implementasi <i>User Interface</i>	35
4.1.3	Implementasi Pemodelan Agen Cerdas dan Lingkungan	41
4.2	Implementasi Kode.....	43
4.2.1	Implementasi penggambaran lintasan.....	43
4.2.2	Implementasi Fungsi Pembelajaran	45
4.2.3	Implementasi Fungsi Pengujian	51
4.3	Evaluasi Sistem	53
4.3.1	Analisis Pembelajaran Berdasarkan Variasi Bentuk Lintasan	54
4.3.2	Analisis Pembelajaran Berdasarkan Lama Waktu Pembelajaran ...	63
4.3.3	Analisis Bentuk Lintasan dan Kemampuan Aksi Agen Cerdas.....	66
BAB V.....		68
KESIMPULAN DAN SARAN.....		68
5.1	Kesimpulan.....	68
5.2	Saran.....	68
DAFTAR PUSTAKA		70
LAMPIRAN.....		

DAFTAR GAMBAR

Gambar 2. 1 Perbedaan hasil agen yang belajar dengan lawan yang berbeda kemampuan	6
Gambar 3. 1 Use Case Diagram Sistem	12
Gambar 3. 2 Flowchart Garis Besar Sistem	17
Gambar 3. 3 Diagram alir fungsi pembuatan lintasan.....	18
Gambar 3. 4 Diagram alir pembelajaran lintasan balap.....	20
Gambar 3. 5 Gambar visualisasi persepsi state agen cerdas	21
Gambar 3. 6 Diagram alir fungsi persepsi state	22
Gambar 3. 7 Diagram alir fungsi Q-learning	24
Gambar 3. 8 Contoh perhitungan 1	25
Gambar 3. 9 Visualisasi ruang arah agen cerdas terhadap arah lintasan.....	27
Gambar 3. 10 Diagram alir aplikasi pengujian agen cerdas.....	29
Gambar 3. 11 Visualisasi bentuk agen cerdas	29
Gambar 3. 12 Contoh rancangan lintasan	30
Gambar 3. 13 Rancangan antar muka aplikasi pembelajaran lintasan.....	32
Gambar 3. 14 Rancangan antar muka aplikasi pengujian agen cerdas	33
Gambar 3. 15 Rancangan aplikasi pembuatan lintasan.....	34
Gambar 4. 1 Halaman Home Sistem.....	36
Gambar 4. 2 Penggambaran lintasan.....	37
Gambar 4. 3 Halaman Pembelajaran Bagian 1	39
Gambar 4. 4 Halaman Pembelajaran Bagian 2	39
Gambar 4. 5 Tampilan Pembaharuan Q-table.....	40
Gambar 4. 6 Contoh proses pengujian agen.....	41
Gambar 4. 7 Potongan kode untuk saat event klik mouse pengguna dilakukan ...	44
Gambar 4. 8 Proses Penyimpanan Lintasan Kedalam Berkas .txt	45
Gambar 4. 9 Badan Fungsi newEpisode	46
Gambar 4. 10 Implementasi Kode Fungsi perceptState.....	47
Gambar 4. 11 Badan utama fungsi drawFrame.....	48

Gambar 4. 12 Potongan Kode Fungsi <code>doLearningTick</code>	49
Gambar 4. 13 Potongan Kode Fungsi <code>Q-learning</code>	49
Gambar 4. 14 Potongan Kode Fungsi <code>getReward</code>	50
Gambar 4. 15 Potongan Kode Inisialisasi <code>R-table</code>	50
Gambar 4. 16 Potongan Kode Muat Memori Pembelajaran	52
Gambar 4. 17 Potongan Kode Fungsi <code>startTesting</code>	52
Gambar 4. 18 Potongan kode fungsi <code>getNextAction</code>	52
Gambar 4. 19 Potongan Kode Fungsi <code>drawFrame</code> Pengujian	53
Gambar 4. 20 Bentuk lintasan sempit	54
Gambar 4. 21 Bentuk lintasan berbalik.....	54
Gambar 4. 22 Bentuk lintasan bundar.....	55
Gambar 4. 23 Bentuk lintasan kombinasi	55
Gambar 4. 24 Grafik peningkatan jumlah keberhasilan agen yang diuji pada lintasan yang tidak dipelajari	63
Gambar 4. 25 Grafik Keberhasilan Pengujian Agen.....	65
Gambar 4. 26 Potongan Data Memori Agen Pada State 71	67
Gambar 4. 27 Potongan Animasi Agen 2 Pada Kasus State 71	67

DAFTAR TABEL

Tabel 3. 1	13
Tabel 3. 2	14
Tabel 3. 3	15
Tabel 3. 4	16
Tabel 3. 5	23
Tabel 3. 6	26
Tabel 3. 7	27
Tabel 4. 1	56
Tabel 4. 2	57
Tabel 4. 3	58
Tabel 4. 4	59
Tabel 4. 5	60
Tabel 4. 6	61
Tabel 4. 7	62
Tabel 4. 8	64
Tabel 4. 9	65
Tabel 4. 10	66
Tabel 4. 11	66

INTISARI

Implementasi Reinforcement Learning Untuk Menentukan Tingkat Kecerdasan Non-Player-Character Pada Permainan Balap Mobil

Q-learning adalah salah satu jenis pembelajaran mesin yang dilakukan pada lingkungan yang tidak diketahui oleh agen, sehingga agen harus mencoba-coba aksi yang tepat pada suatu *state*. Evaluasi aksi yang dilakukan agen adalah memberikan *reward*. Dalam algoritma ini, agen memperbaharui *Q-table* dengan nilai *reward* tersebut. Tabel tersebut menyimpan nilai action untuk setiap *state* selama pelatihan. Nilai tersebut merepresentasikan bobot kesesuaian aksi pada *state* tertentu.

Dengan mengimplementasikan algoritma *Q-learning*, penulis bermaksud mengembangkan sistem dimana terdapat sejumlah NPC pembalap mobil (selanjutnya dikenal sebagai agen) yang masing-masing memiliki kemampuan untuk mempersepsikan kondisi agen pada lingkungan pembelajaran. Aksi yang dapat dilakukan agen untuk menanggapi *state* yang dipersepsikan adalah berbelok ke kiri, maju lurus dan berbelok ke kanan. Tujuan penelitian ini adalah untuk membuktikan bahwa dengan algoritma *Q-learning*, perbedaan tingkat kecerdasan agen cerdas dapat dibuat dengan cara membedakan variasi lintasan yang dipelajari oleh agen dan membedakan variasi waktu belajar agen. Agen akan terbukti lebih pandai apabila semakin banyak variasi lintasan yang dapat diselesaikan saat pengujian.

Dengan menguji agen cerdas pada lintasan uji dengan variasi berbeda, dapat disimpulkan bahwa dapat ditentukan perbedaan tingkat kecerdasan agen dengan cara membedakan jumlah variasi lintasan yang dipelajari oleh agen cerdas dan cara kedua adalah dengan membedakan lama waktu pembelajaran.

Kata kunci : QLearning, Reinforcement Learning, Mobil Cerdas

BAB I

PENDAHULUAN

1.1 Latar Belakang Masalah

NPC (Non-Player-Character) pada aplikasi permainan video merupakan pemain yang tidak dikendalikan oleh manusia. Dalam permainan video, NPC diprogram untuk melakukan tujuan tertentu sehingga terlihat seolah dikendalikan oleh manusia. Untuk melakukan tujuan tersebut, NPC diprogram menggunakan suatu kecerdasan buatan. NPC yang baik dapat menyelesaikan tujuannya dengan efisien dan interaktif. Dalam permainan *single-player*, biasanya NPC dibagi dalam beberapa tingkat kecerdasan/kesulitan (mudah, menengah, atau susah).

Dalam balap mobil, salah satu indikator bahwa pembalap dikatakan lebih baik adalah jika pembalap dapat menyelesaikan banyak variasi lintasan. Dalam permainan video balap mobil *single-player*, NPC pembalap yang berulang kali melintasi sirkuit susah seharusnya dapat melintasi sirkuit yang lebih mudah. Sebelum mencoba sirkuit, NPC sama sekali tidak memiliki pengalaman balap sama sekali. Untuk mencapai tujuannya melintasi garis akhir sirkuit, NPC harus belajar dengan mencoba-coba tindakan tertentu (belok kiri, belok kanan atau maju lurus) pada sebuah kondisi yang dialaminya di lintasan. Saat mencoba-coba lintasan, NPC akan memperbaharui nilai yang digunakan untuk pengambilan aksi yang akan dilakukan, sehingga saat menemui kasus yang sama NPC tidak akan melakukan aksi yang salah. Untuk itu NPC harus mencoba lintasan berulang kali agar dapat mahir pada suatu sirkuit yang sama.

Algoritma pencarian rute tercepat seperti A*, Dijkstra dapat digunakan untuk menyelesaikan permasalahan ini, namun kekurangan algoritma tersebut adalah saat terdapat perubahan aturan dan peta yang lebih kompleks pada permainan. Dimana pengembang harus memodifikasi algoritma agar cocok diterapkan pada aturan permainan yang baru. Selain itu, aksi NPC menjadi lebih

mudah ditebak oleh pemain manusia yang berulang kali memainkan level yang sama.

Reinforcement learning adalah salah satu jenis pembelajaran mesin yang dilakukan pada lingkungan yang tidak diketahui oleh agen, sehingga agen harus mencoba-coba setiap aksi dan akan memperoleh *reward* dari aksi yang dilakukannya. Reward tersebut akan digunakan agen untuk mempertimbangkan aksi pada kondisi sama yang dialaminya pada masa datang. Salah satu algoritma reinforcement learning adalah *Q-Learning*. Dalam algoritma ini, agen memperbaharui table Q, struktur data yang menyimpan nilai pasangan action-state selama pelatihan. Nilai tersebut merepresentasikan reward yang didapatkan agen saat melakukan suatu aksi pada suatu kondisi. Pembaharuan table Q dilakukan selama pelatihan dilakukan. Pembaharuan dilakukan sampai table Q convergence, artinya tidak terdapat perubahan signifikan pada table Q saat ini dengan table Q sebelumnya.

Dengan mengimplementasikan algoritma *Q-Learning*, penulis bermaksud mengembangkan sistem dimana terdapat sejumlah NPC pembalap mobil yang masing-masing dilatih pada sejumlah sirkuit dengan tingkat kesulitan berbeda. Kesulitan sirkuit ditentukan dengan variasi belokan dan variasi objek penghalang di dalamnya. Dari hasil pelatihan tersebut, penulis bermaksud untuk meneliti perbedaan tingkat kesulitan/kecerdasan masing-masing NPC pembalap dengan cara menguji masing-masing NPC pembalap pada suatu sirkuit uji. Cara pengukuran tingkat kecerdasan/kesulitan NPC pembalap adalah dengan menguji seberapa banyak variasi lintasan yang dapat diselesaikan oleh NPC.

1.2 Perumusan Masalah

Berdasarkan latar belakang rumusan masalah yang akan diteliti adalah sebagai berikut :

1. Bagaimana variasi lintasan menentukan tingkat kecerdasan NPC pembalap ?
2. Bagaimana pengaruh lama waktu pembelajaran terhadap kemampuan agen cerdas yang dihasilkan ?

3. Bagaimana korelasi lama waktu pembelajaran dengan waktu tempuh lintasan saat pengujian?

1.3 Batasan Masalah

Berikut ini merupakan batasan sistem untuk penelitian yang akan dilakukan dalam Tugas Akhir :

1. Simulasi aplikasi permainan balap mobil yang akan menjadi dunia permainan adalah aplikasi permainan 2 dimensi dengan sudut pandang kamera dari atas lintasan. Aplikasi akan dikembangkan menggunakan HTML5 *Canvas* dan Javascript dan dibuka dengan browser yang mendukung fitur HTML5 *Canvas*.
2. Kemampuan NPC pembalap dalam permainan hanya sebatas melaju lurus, belok kiri sebesar 15° dan belok kanan sebesar 15° . Apabila NPC pembalap menabrak lintasan atau objek lain, secara otomatis Agen NPC akan dipindahlokasikan ke tengah lintasan.
3. Dalam penelitian ini, implementasi algoritma Q-learning menggunakan konstanta learning rate sebesar 0,2 dan discount factor sebesar 0,8.

1.4 Tujuan Penelitian

Tujuan penelitian tugas akhir ini adalah untuk membuktikan bahwa dengan Q-Learning, NPC pembalap mobil dapat meningkatkan kecerdasan dengan mempelajari lingkungan pelatihan (sirkuit latih) yang semakin sulit.

1.5 Metode Penelitian

Tahapan dalam penelitian ini adalah:

A. Studi Pustaka dan Literatur

Peneliti melakukan pencarian terhadap sumber-sumber yang berkaitan dengan topik penelitian. Topik-topik yang berkaitan dengan algoritma Q-

Learning, *HTML 5* dan *Two Line Segment Interaction Test* akan dipelajari untuk membantu peneliti dalam melakukan proses penelitian.

B. Perancangan Sistem

Tahap perancangan sistem tersebut meliputi perancangan antarmuka, *Use Case diagram* dan diagram alir. Sistem ini akan menerapkan algoritma-algoritma yang mendukung penelitian. Karena sistem ini akan menghasilkan data yang dibutuhkan dalam analisis penelitian, sistem ini akan dirancang sesuai dengan konteks penelitian.

C. Pembangunan Sistem

Pembangunan sistem akan dilakukan pada platform *HTML 5 Canvas* + *Javascript*. Sistem akan dibangun sesuai dengan hasil perancangan sistem dengan mengimplementasikan algoritma yang akan diteliti.

D. Pengumpulan Data Pembelajaran

Setelah sistem selesai dibuat, peneliti akan melakukan pembelajaran terhadap sejumlah agen cerdas untuk disimpan memorinya. Pembelajaran dilakukan untuk pada beberapa variasi lintasan dan beberapa variasi waktu belajar.

E. Pengujian dan Analisis Data

Setelah data memori agen terkumpul, maka selanjutnya dilakukan pengujian terhadap memori yang didapatkan pada lintasan uji. Hasil pengujian adalah selesai atau tidak selesainya pengujian. Setelah pengujian dilakukan, maka hasil pengujian akan dianalisis untuk diambil kesimpulan.

1.6 Sistematika Penulisan

Laporan penelitian untuk tugas akhir ini terdiri dari lima (5) bab. Secara garis besar lima bab tersebut adalah pendahuluan, tinjauan pustaka, analisis dan perancangan sistem, implementasi dan analisis sistem, dan kesimpulan. Penjelasan detail mengenai bab-bab tersebut akan disajikan dalam paragraph di bawah ini.

Bab I, Pendahuluan berisi tentang latar belakang dan gambaran penelitian yang akan dilakukan dalam tugas akhir ini. Deskripsi, alasan, tujuan, dan batasan penelitian disajikan dalam bab ini. Bab II adalah Tinjauan Pustaka yang berisi uraian dan landasan teori yang dipakai dalam penelitian ini. Dalam tinjauan pustaka akan diuraikan mengenai proses seam carving dan deskripsi mengenai algoritma yang dipakai dalam implementasi sistem untuk penelitian ini.

Bab III dan bab IV akan membahas mengenai sistem dalam penelitian ini. Analisis kebutuhan sistem dan perancangan sistem akan dibahas dalam Bab III. Bab IV membahas implementasi sistem dan analisis data. Segala hal yang berkaitan dengan sistem dan analisisnya disajikan dalam kedua bab tersebut. Bab III dan bab IV inilah yang nantinya akan dijadikan landasan dalam bab V yaitu kesimpulan.

Bab V, Kesimpulan berisi jawaban dari rumusan masalah yang terdapat dalam Bab I. Jawaban yang didapatkan adalah hasil dari penelitian yang dijabarkan dalam Bab III dan IV. Dalam bab ini juga akan ditulis saran untuk penelitian selanjutnya.

BAB V

KESIMPULAN DAN SARAN

5.1 Kesimpulan

Berdasarkan hasil implementasi dan evaluasi sistem, diperoleh kesimpulan bahwa dengan menggunakan algoritma *Q-learning* dapat ditentukan perbedaan tingkat kecerdasan agen. Beberapa poin yang dapat digunakan untuk memperjelas keberhasilan tersebut adalah sebagai berikut.

1. Seiring bertambahnya variasi lintasan yang dipelajari agen cerdas pada suatu tahap, semakin banyak agen yang menyelesaikan lintasan uji yang tidak dipelajari pada tahap tersebut.
2. Pembelajaran agen pada lintasan kombinasi menunjukkan bahwa semakin lamanya waktu pembelajaran yang dilakukan oleh agen cerdas, semakin banyak rata-rata agen cerdas yang dapat menyelesaikan lintasan kombinasi dan lintasan variasi (putar balik, sempit dan bundar).
3. Tidak ada korelasi antara lamanya waktu pembelajaran dengan lamanya waktu untuk menyelesaikan lintasan uji.

5.2 Saran

Saran untuk pengembangan dan perbaikan sistem adalah dengan beberapa poin berikut

1. Penelitian terkait besarnya sudut belokan, panjangnya sensor pendeteksi objek dan pembuatan lintasan yang sesuai dengan aksi agen cerdas perlu dilakukan agar tidak terjadi *bug* saat pembelajaran dilakukan.
2. Distribusi mulai pembelajaran pada setiap segmen arah sebaiknya dilakukan agar pembelajaran terjadi secara merata untuk setiap segmennya. Pada penelitian ini, setiap agen yang menabrak akan

kembali ke garis start, sehingga waktu untuk mempelajari state yang dapat dipersepsi pada segmen awal lebih lama daripada segmen-segmen lain. Dengan mendistribusikan waktu secara merata untuk tiap segmen yang dilalui agen cerdas, diharapkan pembelajaran state yang dapat dipersepsi pada tiap segmen dapat terjadi secara merata.

©UKDW

DAFTAR PUSTAKA

- H. Cormen, T., E. Leiserson, C., & Stein, C. (2001). Introduction to Algorithms, Second Edition. *The MIT Press* , Cambridge, MA.
- Karavolos, D. (2013). Q-learning with heuristic exploration in Simulated Car Racing. *Theses* , University of Amsterdam.
- Khriji, L., Touati, F., Benhmed, K., & Al-Yahmedi, A. (2011). Mobile robot Navigation Based on Q-Learning Technique. *International Journal of Advanced Robotic Systems* , 45-51.
- Massoudi, P., & Fassihi, A. (2013). Achieving Dynamic AI Difficulty by Using Reinforcement Learning and Fuzzy Logic Skill Metering. *IEEE International Games Innovation Conference* , 163-168.
- Patel, P. (2009). Improving Computer Game Bots behavior using Q-Learning. *OpenSIUC Theses* , 76.
- Sedgewick, R. (1983). ALGORITHMS. *Addison-Wesley Publishing Company* .