

Dr. Ir. Agus Wibowo, M.Kom, M.Si, MM

REGULASI APLIKASI AI (Artificial Intelligence)



YAYASAN PRIMA AGUS TEKNIK



Dr. Ir. Agus Wibowo, M.Kom, M.Si, MM

REGULASI APLIKASI AI

(Artificial Intelligence)



YAYASAN PRIMA AGUS TEKNIK

PENERBIT :

YAYASAN PRIMA AGUS TEKNIK
Jl. Majapahit No. 605 Semarang
Telp. (024) 6723456. Fax. 024-6710144
Email : penerbit_ypat@stekom.ac.id

ISBN 978-634-7227-46-1 (PDF)



9

786347

227461

REGULASI APLIKASI AI (Artificial Intelligence)

Penulis :

Dr. Ir. Agus Wibowo, M.Kom, M.Si, MM

ISBN : 978-634-7227-46-1 (PDF)

Editor :

Dr. Joseph Teguh Santoso, S.Kom., M.Kom.

Penyunting :

Dr. Mars Caroline Wibowo. S.T., M.Mm.Tech

Desain Sampul dan Tata Letak :

Irdha Yuniarto, S.Ds., M.Kom

Penebit :

Yayasan Prima Agus Teknik Bekerja sama dengan
Universitas Sains & Teknologi Komputer (Universitas STEKOM)

Anggota IKAPI No: 279 / ALB / JTE / 2023

Redaksi :

Jl. Majapahit no 605 Semarang

Telp. 08122925000

Fax. 024-6710144

Email : penerbit_ypat@stekom.ac.id

Distributor Tunggal :

Universitas STEKOM

Jl. Majapahit no 605 Semarang

Telp. 08122925000

Fax. 024-6710144

Email : info@stekom.ac.id

Hak cipta dilindungi undang-undang

Dilarang memperbanyak karya tulis ini dalam bentuk dan dengan cara
apapun tanpa ijin dari penulis

KATA PENGANTAR

Puji syukur kami panjatkan ke hadirat Tuhan Yang Maha Esa atas terselesaikannya buku berjudul *Regulasi Aplikasi AI* ini hadir sebagai kajian komprehensif mengenai aspek hukum, sosial, dan teknis yang mengiringi perkembangan pesat kecerdasan buatan (AI) di berbagai bidang. Setiap bab disusun untuk memberikan pemahaman mendalam dan analisis kritis terhadap regulasi yang mengatur pemanfaatan AI, mulai dari perlindungan data pribadi hingga kebijakan militer yang memengaruhi hak-hak sipil.

Bab pertama membahas kerangka regulasi kecerdasan buatan dan teknologi informasi di Uni Eropa, menyoroti aspek transparansi, keterjelasan, akuntabilitas, serta investasi dan koordinasi regulasi di masa depan, termasuk jalur legislatif dan peran pengadilan. Bab kedua mengangkat dampak pengenalan teknologi pengenalan wajah (FRT) berbasis AI terhadap privasi individu, mengulas berbagai tantangan hukum dan perlindungan data di Indonesia.

Pada bab ketiga, fokus beralih pada peran AI dalam dinamika politik dan psikologi pemerintahan, mengeksplorasi ancaman psikologis baru, klasifikasi muai, serta respons hukum di berbagai negara seperti Uni Eropa, Amerika Serikat, Tiongkok, dan Rusia. Bab keempat mengkaji pemanfaatan AI dalam investasi kewarganegaraan, mencakup aspek uji tuntas, kepatuhan korporasi dan imigrasi, serta tantangan yang muncul dengan pengenalan AI. Bab kelima menguraikan potensi dan tantangan AI di bidang layanan kesehatan, mulai dari keputusan perawatan hingga pengaturan otonomi dan hubungan dokter-pasien, serta dampak regulasi terhadap sistem kesehatan Indonesia. Bab keenam memfokuskan pada praktik hukum yang didukung AI sekaligus membahas regulasi yang mengaturnya.

Bab ketujuh meneliti perlindungan hak cipta dalam musik digital yang dihasilkan AI, menjelaskan evolusi musik, legalitas, kreativitas AI, serta solusi alternatif atas tantangan hak cipta karya digital. Bab kedelapan mengulas manajemen risiko penerbangan dengan AI, meninjau situasi terkini, peran AI dalam keselamatan, dan respons hukum global terhadap kerangka regulasi.

Bab kesembilan membahas AI otonom dalam pelabuhan pintar dan manajemen rantai pasok, membicarakan integrasi sistem, potensi risiko, dan intervensi regulasi yang optimal. Bab kesepuluh menyajikan AI dalam kebijakan iklim dan energi, khususnya di UE dan Jepang, meliputi insentif, kesepakatan global, dan model regulasi cerdas.

Bab kesebelas hingga ketiga belas mengkaji regulasi AI di bidang militer dan konflik bersenjata, termasuk senjata otonom, kewajiban perlindungan sipil, tanggung jawab negara atas pelanggaran hukum perang, serta kendali manusia atas senjata otonom, dengan studi kasus kebijakan pertahanan Amerika Serikat.

Melalui pemaparan sistematis ini, buku ini bertujuan memberikan wawasan holistik bagi akademisi, praktisi hukum, pembuat kebijakan, dan seluruh pemangku kepentingan terkait untuk memahami dan mengelola regulasi AI di era digital yang penuh tantangan dan peluang.

Semarang, September 2025

Penulis

Dr. Ir. Agus Wibowo, M.Kom, M.Si, MM

DAFTAR ISI

KATA PENGANTAR.....	ii
DAFTAR ISI.....	iii
BAB 1 KECERDASAN BUATAN DAN REGULASI ITX DI UNI EROPA.....	1
1.1 Pendahuluan.....	1
1.2 Transparansi.....	2
1.3 Keterangan.....	4
1.4 Akuntabilitas.....	6
1.5 Investasi Dan Koordinasi.....	8
1.6 Cakupan Regulasi Di Masa Depan?	9
1.7 Jalur Legislatif	9
1.8 Pengadilan	10
BAB 2 DAMPAK PENGENALAN WAJAH AI TERHADAP PRIVASI.....	14
2.1 Konsep Dan Fungsi FRT.....	14
2.2 Bagaimana Ai Memengaruhi FRT?	14
2.3 Apakah Citra Wajah Merupakan Data Pribadi?	15
2.4 Ai Dan Privasi: Pengenalan Wajah Di Indonesia.....	16
2.5 Tantangan Hukum Terkait FRT	17
2.6 Privasi Dan Perlindungan Data	19
BAB 3 AI, PEMERINTAH, POLITIK, DAN PSIKOLOGI	24
3.1 Pendahuluan.....	24
3.2 Pengaruh Ai Dalam Dinamika Politik Indonesia	24
3.3 Ancaman Psikologis Baru Terhadap Pemerintah Dan Institusi Politik	25
3.4 Klasifikasi Muai	27
3.5 Tanggapan Hukum Atas Muai Di Uni Eropa Dan AS	27
3.6 Tanggapan Hukum Di Tiongkok Dan Rusia	31
3.7 Pendekatan Sistemik Bagi Keberhasilan Hukum Balasan.....	34
3.8 Kesimpulan	35
BAB 4 PEMANFAATAN AI UNTUK KEWARGANEGARAAN MELALUI INVESTASI	37
4.1 Pendahuluan.....	37
4.2 Uji Tuntas	38
4.3 Kepatuhan Korporasi Dalam CBI.....	40
4.4 Kepatuhan Imigrasi Dalam CBI	42
4.5 Challenge Dengan Pengenalan Ai.....	43
4.6 Kesimpulan	45
BAB 5 POTENSI, TANTANGAN, REGULASI AI BIDANG KESEHATAN.....	46
5.1 Ai Dalam Layanan Kesehatan	46
5.2 Potensi Ai Dalam Keputusan Perawatan Kesehatan	47
5.3 Tantangan, Resiko, Keandalan, Dan Terpercaya	50

5.4	Tantangan Otonomi Dan Relasi Dokter-Pasien.....	53
5.5	Pengaturan Keputusan Dalam Layanan Kesehatan Berbasis AI	54
5.6	Dampak AI Pada Sistem Pelayanan Kesehatan Di Indonesia.....	59
5.7	Kesimpulan	59
BAB 6	KECERDASAN BUATAN DALAM PRAKTIK HUKUM.....	61
6.1	Kemampuan AI Dalam Praktik Hukum Saat Ini	61
6.2	Regulasi.....	65
6.3	Kesimpulan	69
BAB 7	AI DALAM PERLINDUNGAN HAK CIPTA MUSIK DIGITAL	70
7.1	Evolusi Ke Revolusi Musik.....	70
7.2	Legalitas Musik Berbasis Kecerdasan Buatan Di Indonesia.....	72
7.3	Kemampuan AI Untuk Kreativitas.....	73
7.4	Tantangan Hak Cipta Musik Buatan AI.....	75
7.5	Solusi Alternatif Karya Yang Dihasilkan Komputer	78
7.6	Kesimpulan	81
BAB 8	AI DALAM MANAJEMEN RISIKO PENERBANGAN	83
8.1	Pendahuluan.....	83
8.2	Keadaan AI Saat Ini Dalam Penerbangan	83
8.3	Peran AI Dalam Peningkatan Keselamatan Penerbangan Di Indonesia	86
8.4	AI Dan Potensi Manajemen Risiko Penerbangan	87
8.5	Respons Global Atas Kerangka Hukum.....	89
8.6	Kesimpulan	91
BAB 9	AI OTONOM, PELABUHAN PINTAR, DAN MANAJEMEN RANTAI PASOK	93
9.1	AI, Pelabuhan Pintar, Dan Kapal Cerdas	93
9.2	Kecerdasan Buatan Pelabuhan Pintar Dan CRM Di Indonesia	94
9.2	Integrasi Dan Kolaborasi.....	95
9.3	Potensi Risiko Dan Masalah Regulasi	96
9.4	Menuju Intervensi Regulasi Yang Optimal.....	97
9.5	Kesimpulan	99
BAB 10	AI DALAM KEBIJAKAN IKLIM DAN ENERGI UE–JEPANG	101
10.1	Pendahuluan.....	101
10.2	Insentif Dan Kesepakatan Global.....	102
10.3	AI Dan Kebijakan Iklim-Energi Jepang	105
10.4	Model Regulasi Cerdas	108
10.5	Kesimpulan	109
BAB 11	REGULASI AI MILITER UNTUK PERLINDUNGAN SIPIL.....	111
11.1	Pendahuluan.....	111
11.2	Otonomi Dan Sarana Perang	111
11.3	Senjata Otonom Dan Perdebatan GGE Tentang Laws	116
11.4	Kewajiban Perlindungan	117
11.5	Kesimpulan	123

BAB 12 IMPLIKASI HUKUM AI DALAM KONFLIK BERSENJATA	125
12.1 Pendahuluan.....	125
12.2 Tanggung Jawab Negara Atas Pelanggaran Hukum Perang.....	125
12.3 Rezim Pertanggungjawaban	131
12.4 Kesimpulan	133
BAB 13 KENDALI MANUSIA ATAS SENJATA OTONOM.....	134
13.1 Pendahuluan.....	134
13.2 Masalah Dalam Kebijakan Pertahanan AS.....	134
13.3 Asumsi Kebijakan Pertahanan AS Terkait AWS.....	138
13.4 Representasi Pertahanan AS Yang Luput.....	146
13.5 Kesimpulan	151
13.6 Ringkasan.....	152
DAFTAR PUSTAKA	154

BAB 1

KECERDASAN BUATAN DAN REGULASI ITX DI UNI EROPA

Terdapat pengakuan bahwa Uni Eropa (UE) tertinggal dari para pesaing utama dalam persaingan untuk menarik, menciptakan, dan membina perusahaan serta investasi AI. Atas dasar ini, Komisi Eropa telah menerbitkan Buku Putih, dan proses konsultasi publik mengenai hal ini kini telah dilakukan. Rencananya, proposal legislatif akan segera diterbitkan.

Dengan adanya beberapa ketidakpastian mengenai regulasi AI yang tepat di masa mendatang di UE, bab ini tetap berusaha untuk menetapkan dari Buku Putih dan proses konsultasi indikasi apa yang ada bagi perusahaan terkait poin-poin utama yang perlu diperhatikan: transparansi, kemudahan dijelaskan, dan akuntabilitas. Bab ini mengontekstualisasikan konsep-konsep ini dan menyoroti tantangan yang melekat di dalamnya. Tentu saja, ini merupakan tantangan yang umum bagi semua yurisdiksi.

Namun, masih terdapat masalah yang lebih unik bagi Uni Eropa, terutama skeptisisme mendasar mengenai AI, tidak hanya di kalangan warga negara Uni Eropa, tetapi juga otoritas publik dan bisnis—dengan 90% responden konsultasi menyatakan kekhawatiran bahwa AI dapat melanggar hak-hak fundamental. Terdapat pula perkembangan menarik, yaitu Parlemen Eropa telah mempublikasikan posisinya mengenai AI dan munculnya nasionalisme, baik dalam konsultasi maupun setelah pandemi COVID-19, yang memberikan gambaran yang jauh dari seragam.

Kenyataannya, regulasi di lapangan kemungkinan masih akan memakan waktu bertahun-tahun lagi.

Mengingat tindakan melalui legislasi ('harmonisasi positif') masih jauh dari mendesak, penting untuk menyoroti kemungkinan peran Pengadilan Keadilan Uni Eropa dalam mengembangkan bidang ini melalui yurisprudensinya (yang disebut 'harmonisasi negatif'). Secara keseluruhan, perusahaan akan terdorong oleh semangat Pengadilan, meskipun nada spesifiknya sulit diidentifikasi pada tahap ini.

Pada akhirnya, terkait masa depan regulasi AI di Uni Eropa, harus diakui bahwa meskipun terdapat ambisi untuk mencapai banyak hal di bidang ini, visi dan filosofi kolektif tentang apa yang dapat dilakukan AI, dan bagaimana AI dapat bermanfaat bagi masyarakat, masih terfragmentasi.

1.1 PENDAHULUAN

Buku Putih mendefinisikan AI sebagai 'kumpulan teknologi yang menggabungkan data, algoritma, dan daya komputasi'. Aplikasinya beragam dan Komisi mengakui bahwa, secara positif, AI dapat meningkatkan layanan kesehatan, membantu memerangi perubahan iklim, meningkatkan efisiensi produksi, dan meningkatkan keamanan (di antara manfaat lainnya). Namun, untuk mencapai ambisi ini, AI perlu mengatasi tantangan tertentu, termasuk fakta bahwa warga negara merasa cemas mengenai kapasitas AI untuk merugikan mereka, baik

secara sengaja maupun tidak sengaja, dan bahwa bisnis sedang mencari kepastian hukum. Sebuah cara terpadu untuk maju sedang diupayakan.

Bagaimana ini dapat dicapai? Jawabannya tidak langsung dan bergantung pada sejumlah faktor subjektif yang berakar pada kepercayaan dan etika. Bagi Uni Eropa, terdapat kesadaran yang semakin meningkat bahwa membuat keputusan dengan cepat atau lebih efisien bukanlah penentu keberhasilan AI. Yang lebih penting adalah bagaimana manfaatnya diwujudkan. Apakah kita siap menyerahkan keputusan yang berdampak pada kehidupan manusia ke tangan mesin? Mungkin jawaban atas 'kesiapan' ini terletak di koridor kepercayaan dan etika, dan kepada konsep-konsep inilah bab ini sekarang beralih.

Pada bulan April 2019, Kelompok Pakar Tingkat Tinggi Komisi Eropa tentang AI mengadopsi Pedoman Etika untuk AI yang Dapat Dipercaya, yang menekankan bahwa manusia hanya akan dapat dengan yakin dan sepenuhnya menuai manfaat AI jika mereka dapat mempercayai teknologinya. Prinsip umum bahwa inisiatif AI tidak boleh diwujudkan jika memerlukan kompromi terhadap etika sudah sangat jelas. Namun, kepatuhan terhadap etika merupakan kompleksitas yang berbeda yang patut dibahas. Interpretasi 'kreatif' atas prinsip-prinsip etika tidak boleh digunakan sebagai tameng untuk mencapai kepatuhan 'semata-mata' di mana perusahaan *tampaknya hanya* mematuhi.

Hubungan antara etika dan hukum menambah lapisan kerumitan lebih lanjut. Sebagaimana disoroti oleh Pengawas Perlindungan Data Eropa, etika di Uni Eropa tidak dipahami sebagai alternatif kepatuhan terhadap hukum, melainkan sebagai nilai-nilai dasar yang melindungi martabat dan kebebasan manusia.

Bab ini akan mencoba mengurai kompleksitas di balik tiga prinsip etika yang membentuk landasan model AI yang efektif—transparansi, kemudahan dijelaskan, dan akuntabilitas.

1.2 TRANSPARANSI

Dalam konteks AI, transparansi menunjukkan kemampuan untuk mendeskripsikan, memeriksa, dan mereproduksi mekanisme yang digunakan sistem AI dalam mengambil keputusan. Namun, transparansi berkaitan erat dengan kepercayaan, yang juga berarti bersikap eksplisit dan terbuka tentang pilihan dan keputusan terkait sumber data, proses pengembangan, dan pemangku kepentingan. Intinya, transparansi berarti memiliki pandangan yang utuh terhadap sistem pada tiga tingkatan. Pertama, pada tingkat implementasi, di mana model AI bertindak berdasarkan data yang dimasukkan untuk menghasilkan keluaran yang diketahui, termasuk prinsip-prinsip teknis model dan parameter terkait. Komisi menganggap ini sebagai model 'kotak putih' standar, berbeda dengan model 'kotak hitam' di mana prinsip ini tidak diketahui. Kedua, pada tingkat spesifikasi, semua informasi yang menghasilkan implementasi, seperti tujuan, tugas, dan kumpulan data yang relevan, harus terbuka dan diketahui. Tingkat ketiga adalah interpretabilitas, yang merepresentasikan pemahaman tentang mekanisme yang mendasari model. Misalnya, apa prinsip logis di balik pemrosesan data dan apa alasan di balik keluaran tertentu? Komisi berpendapat bahwa pertanyaan-pertanyaan ini adalah yang paling sulit dijawab dan tingkat transparansi ketiga tidak tercapai

dalam sistem AI saat ini. Hal ini khususnya rumit karena tingkat transparansi ketiga terkait erat dengan keadilan dalam pengambilan keputusan, yang seringkali berdampak pada kehidupan manusia. Keadilan itu sendiri merupakan konsep subjektif dan kontekstual, yang dipengaruhi oleh berbagai faktor sosial, budaya, dan hukum.

Dengan tingkat subjektivitas yang ada, penting untuk diingat bahwa model AI tidak dirancang oleh mesin dan berpotensi mencerminkan bias dan prasangka perancang yang memilih fitur, metrik, dan struktur model. Konsep kepercayaan juga perlu diklarifikasi lebih lanjut untuk memahami efektivitas AI. Kepercayaan terutama digunakan untuk berbicara tentang kepercayaan atau ketidakpercayaan individu dan lembaga yang bertanggung jawab untuk mengembangkan, menerapkan, atau menggunakan AI. Namun, kepercayaan juga dapat diarahkan kepada mereka yang bertanggung jawab untuk mengatur AI, yang mengawasi penggunaannya atau menyediakan kondisi untuk pengembangannya. Melibatkan semua pemangku kepentingan ini dalam pengambilan keputusan yang berdampak pada manusia mungkin merupakan cara paling adil untuk membangun kepercayaan, tetapi hal ini mungkin tidak memungkinkan karena jumlah pemangku kepentingan atau kompleksitas sistem AI. Hasilnya kemudian menjadi hampir sama pentingnya dengan prosesnya—tetapi mampukah kita memastikan bahwa sistem AI setransparan mungkin?

Ketidakjelasan dalam pembelajaran mesin, yang disebut algoritma 'kotak hitam', sering disebut sebagai salah satu hambatan utama untuk mencapai transparansi dalam AI. Tujuan utama algoritma AI adalah untuk mengidentifikasi pola atau kesamaan dalam kumpulan data tertentu. Dengan demikian, korelasi yang disajikan oleh kumpulan data tertentu (misalnya, korelasi antara wilayah geografis dan ras) juga dipelajari oleh algoritma. Ada risiko bahwa korelasi ini akan menyebabkan bias dan hasil yang tidak adil karena prasangka manusia. Misalnya, terdapat banyak masalah terkait model pengenalan wajah yang dilatih pada kumpulan data terbatas yang mencerminkan bias pengembang model, dan hasil yang tidak adil dari model keputusan pinjaman yang secara historis menggunakan kumpulan data bias untuk menentukan ketersediaan kredit.

Oleh karena itu, pengambilan keputusan yang bebas bias harus menjadi hasil dari model AI yang transparan. Proses pengambilan keputusan harus terbuka dan adil. Misalnya, aplikasi pinjaman yang ditolak harus dapat mencantumkan alasan objektif mengapa pemohon ditolak. Alasan ini bisa berupa pendapatan tahunan atau jumlah tabungan, tetapi bukan alasan subjektif seperti ras, jenis kelamin, atau kode pos. Kriteria objektif yang transparan juga akan menghasilkan proses yang lebih adil untuk menantang keputusan AI, karena pengguna AI akan berada dalam posisi yang lebih baik untuk memahami keputusan AI yang tidak terduga. Untuk mencapai hasil yang paling adil, harus dimungkinkan untuk menunjukkan bahwa proses desain dan implementasi yang telah dilakukan untuk membuat keputusan tertentu adil secara etis dan dapat dipercaya.

Melabeli model AI sebagai efektif tanpa mencapai tingkat transparansi yang dibahas di atas membawa risiko inheren di tingkat pengguna. Hal ini mendorong 'ketidaktahuan', karena pengguna tidak menyadari tindakan mereka ketika mereka membiarkan AI membuat keputusan. Hal ini dapat menyebabkan ketergantungan yang berlebihan pada AI, yang

mungkin tidak sejalan dengan gagasan moralitas masyarakat.¹⁸ Agar AI berhasil, objek kepercayaan perlu diperjelas, bersamaan dengan mengidentifikasi berbagai hubungan kepercayaan dengan berbagai pihak.

Oleh karena itu, Buku Putih mengakui tantangan dalam memverifikasi kepatuhan, dan mempertimbangkan untuk memberikan tanggung jawab kepada perusahaan untuk secara cermat menyimpan data pelatihan (dan menjelaskan mengapa data tersebut dipilih), mendokumentasikan pemrograman algoritma, dan mencatat proses/teknik yang relevan saat merancang atau menguji sistem. Idenya bukan hanya bahwa hal ini dapat menciptakan tolok ukur untuk menilai kepatuhan, tetapi juga dapat memitigasi risiko dan meningkatkan proses perancangan AI dengan mengintegrasikan praktik-praktik yang lebih baik dalam transparansi sejak awal. Hal ini masih belum mencapai tingkat transparansi ketiga, yang telah dijelaskan sebelumnya.

1.3 KETERJELASAN

Sementara transparansi berkaitan dengan pemahaman tentang bagaimana model AI membuat keputusannya, keterjelasan melangkah lebih jauh. Transparansi menambahkan kebutuhan akan justifikasi, yang dalam konteks AI akan membutuhkan penjelasan mengapa suatu keputusan tertentu diambil. Misalnya, jika permohonan pinjaman ditolak, konsumen berhak menanyakan alasan atau justifikasi di balik penolakan tersebut. Menyelesaikan masalah-masalah ini secara definitif sangatlah mendesak karena terdapat perdebatan akademis yang sedang berlangsung mengenai cakupan persyaratan hukum yang berlaku saat ini berdasarkan Peraturan Perlindungan Data Umum (GDPR).

Keterjelasan juga dikaitkan dengan tanggung jawab, yaitu para ahli AI harus tahu apa yang mereka lakukan, dan harus mampu serta bersedia untuk berkomunikasi, menjelaskan, dan memberikan alasan atas dampak yang ditimbulkan model AI terhadap subjek manusia dan non-manusia. Membaca yang tersirat, moralitas, sekali lagi, terkait erat, termasuk kewajiban untuk memperoleh kesadaran yang lebih besar akan konsekuensi yang tidak diinginkan dan signifikansi moral dari apa yang dilakukan model tersebut. Kesadaran ini telah dipandang sebagai persyaratan penting bagi kerja AI yang efektif.

Keterjelasan dapat mengambil berbagai bentuk, yang paling jelas adalah justifikasi fitur-fitur utama yang berperan dalam keputusan AI. Sebaliknya, menjelaskan persyaratan yang tidak terpenuhi juga dapat membantu dalam menjelaskan suatu keputusan, terutama dalam skenario yang berhadapan langsung dengan nasabah seperti pengajuan pinjaman, yang mungkin menjelaskan mengapa nasabah belum memenuhi kriteria tertentu untuk berhasil. Terlepas dari justifikasi yang diberikan, risiko subjektivitas yang mencerminkan bias manusia masih tetap ada. Komisi berpandangan bahwa bias dan diskriminasi merupakan risiko inheren dari setiap aktivitas sosial atau ekonomi yang melibatkan pengambilan keputusan manusia. Namun, AI memperkuat risiko ini dengan memengaruhi sekelompok besar orang yang mungkin secara langsung terdampak oleh diskriminasi dan bias tersebut. Hal ini khususnya berlaku ketika AI 'belajar' dari data yang dimasukkan ke dalamnya.

Lalu, bagaimana kita memastikan bahwa model AI adil dalam proses pengambilan keputusannya? Jawabannya tidak meyakinkan; namun, GDPR menawarkan sedikit kelonggaran, setidaknya terkait pemrosesan data. GDPR secara khusus diberlakukan untuk memastikan perlindungan data pribadi dalam konteks perkembangan teknologi yang pesat. GDPR menetapkan bahwa data pribadi harus diproses secara adil, yang menyiratkan analisis apakah pemrosesan tersebut diskriminatif, merugikan, atau menyesatkan bagi individu yang terdampak. Ini berarti bahwa pengendali AI harus mempertimbangkan kemungkinan dampak penggunaan AI secara berkala dan terus-menerus menilai ulang, yang membawa kita kembali pada pentingnya keterjelasan.

GDPR menjunjung tinggi prinsip transparansi dan keterjelasan dengan memasukkan hak subjek data atas informasi dan akses ke data pribadi. Ini berarti bahwa pengembang sistem AI memiliki kewajiban hukum untuk membangun perlindungan yang menjamin penegakan prinsip-prinsip perlindungan data. Namun, definisi data pribadi tidak dijelaskan secara mendalam, sehingga sulit untuk menentukan sejauh mana AI dapat digunakan untuk keperluan pemrosesan data. Jika AI ‘mempelajari’ pola pemrosesan data yang berada di luar cakupan GDPR, maka AI berisiko mengalami diskriminasi dan bias, seperti yang telah disebutkan sebelumnya. Hal ini menjadikan penting untuk mengatasi kesenjangan dalam pendefinisian data pribadi pada tahap pengembangan.

Sekilas, larangan GDPR atas pemrosesan kategori khusus data pribadi hanya dengan cara otomatis dapat dilihat memberikan perlindungan yang kuat terhadap diskriminasi oleh model AI. Namun, pada kenyataannya, kategori khusus data pribadi tidak mencakup warna kulit, bahasa, properti, keanggotaan minoritas, dan kelahiran, yang diakui sebagai dasar diskriminasi dalam Piagam Uni Eropa. Hal ini merupakan potensi kesenjangan dalam pencegahan diskriminasi melalui pemrosesan data pribadi. GDPR juga melarang pembuatan profil, yaitu suatu bentuk pemrosesan di mana sistem AI mengkategorikan individu berdasarkan probabilitas mereka untuk termasuk dalam kelompok tertentu, dan bukan sebagai individu yang berdiri sendiri. Namun, persetujuan khusus subjek data merupakan pengecualian terhadap larangan tersebut, yang sekali lagi menyisakan ruang inefisiensi untuk mencegah diskriminasi.

Dalam Buku Putihnya, Komisi mengakui bahwa penting untuk meninjau kerangka legislatif saat ini untuk mencerminkan perkembangan teknologi terkini dan untuk sepenuhnya mempertimbangkan pertimbangan kemanusiaan dan etika. AI. Ketidakjelasan sebagai fitur inheren AI membuat implementasi undang-undang apa pun di sekitarnya menjadi lebih menantang. Dari sudut pandang konsumen, tingkat keamanan dan perlindungan hak-hak fundamental seharusnya tidak berbeda dengan produk yang tidak bergantung pada AI untuk pengembangan atau penerapannya. AI juga merupakan lanskap yang berkembang pesat, dan undang-undang apa pun yang ada saat ini harus cepat beradaptasi dan berevolusi agar kerangka regulasinya efektif.

Proposal-proposal yang muncul saat ini tampaknya kurang memadai terkait dengan keterjelasan (yang memang diharapkan mengingat transparansi penuh juga belum tercapai). Pada akhirnya, terdapat usulan bahwa perusahaan mungkin diwajibkan untuk memberikan

'informasi yang jelas ... mengenai kemampuan dan keterbatasan sistem AI' dan bahwa 'warga negara harus diberi informasi yang jelas ketika mereka berinteraksi dengan sistem AI, bukan manusia'. Namun, hal ini tampaknya hanya solusi sementara untuk mengurangi kekurangan yang diantisipasi dari keterjelasan, alih-alih langkah-langkah yang menyediakannya.

1.4 AKUNTABILITAS

Transparansi, keterjelasan, dan akuntabilitas mungkin sering digunakan sebagai sinonim, tetapi penting untuk menyoroti perbedaan mendasar di antara keduanya. Sementara transparansi dan keterjelasan mengacu pada penjelasan keputusan sebelum model AI mengambil keputusannya, akuntabilitas mengacu pada kemampuan untuk menjelaskan peran model AI setelah tindakan dilakukan, atau tidak dilakukan. Secara sederhana, tujuannya adalah untuk menentukan di mana letak tanggung jawabnya.

Tidak mengherankan jika hanya ada sedikit aturan yang tegas dan jelas tentang tanggung jawab. Namun, dua isu utama perlu dipertimbangkan. Pertama, ada pertanyaan tentang bagaimana tanggung jawab akan didistribusikan di antara berbagai pelaku AI yang terlibat. Banyak aktor yang terlibat dalam siklus hidup sistem AI. Ini termasuk pengembang, deployer (orang yang menggunakan produk atau layanan yang dilengkapi AI) dan kemungkinan pihak lain (produsen, distributor atau importir, penyedia layanan, pengguna profesional atau swasta).

Menurut Komisi, dalam kerangka peraturan di masa mendatang, setiap tanggung jawab harus disesuaikan dengan pihak yang paling tepat untuk mengatasi potensi risiko apa pun. Misalnya, meskipun pengembang AI mungkin paling tepat untuk mengatasi risiko yang timbul dari fase pengembangan, kemampuan mereka untuk mengendalikan risiko selama fase penggunaan mungkin lebih terbatas. Dalam hal ini, deployer harus tunduk pada kewajiban yang relevan. Hal ini tanpa mengurangi pertanyaan, untuk tujuan tanggung jawab kepada pengguna akhir atau pihak lain yang dirugikan dan memastikan akses yang efektif terhadap keadilan, pihak mana yang harus bertanggung jawab atas kerugian yang ditimbulkan. Berdasarkan hukum pertanggungjawaban produk Uni Eropa, pertanggungjawaban atas produk cacat dibebankan kepada produsen, tanpa mengurangi hukum nasional, yang juga memungkinkan pemulihan dari pihak lain. Konsep produk cacat muncul ketika produk tersebut tidak memberikan keamanan yang berhak diharapkan seseorang, dengan mempertimbangkan semua keadaan, termasuk penyajian, penggunaan yang secara wajar dapat diharapkan, dan waktu produk tersebut beredar. Hal ini terbatas penggunaannya dalam konteks AI karena tidak terpadu, dan AI mampu menimbulkan kerugian tanpa menjadi 'cacat'.

Pendekatan Komisi dapat dilihat sebagai apa yang disebut sebagai kondisi 'kendali'. Seorang agen AI seharusnya hanya bertanggung jawab atas suatu tindakan atau keputusan jika mereka memiliki kendali atas tindakan mereka. Namun, bagaimana jika agen tersebut bukan manusia dan tanggung jawabnya terletak pada mesin? Ada perdebatan luas tentang masalah ini tentang apakah teknologi dapat menjadi agen yang bertanggung jawab. Karena sifat model AI melibatkan pemrosesan data dan 'pembelajaran' selama periode waktu tertentu, mengaitkan tanggung jawab penuh kepada manusia mungkin tampak ekstrem. Pada saat yang

sama, tidak dapat diabaikan bahwa manusia berada di balik data yang dimasukkan ke dalam model AI, setidaknya pada awalnya. Sekalipun (beberapa) AI dapat bertindak atau memutuskan, mereka tidak memiliki kapasitas agensi moral, sehingga tanggung jawab atas tindakan mereka tetap berada di tangan agen manusia yang mengembangkan dan menggunakan teknologi tersebut. Untuk mengatasi batasan tipis antara intervensi manusia dan AI, pendekatan hibrida 'agensi terdistribusi' telah disarankan, yang mencakup tanggung jawab terdistribusi di mana semua pihak bertanggung jawab atas peran mereka dalam hasil dan tindakan aplikasi AI berdasarkan kerangka moral yang sama. Namun, konsep ini berisiko dalam implementasinya, di mana suatu pihak dapat dianggap bertanggung jawab secara etis tetapi tidak secara hukum. Sekalipun pendekatan ini diadopsi, kenyataannya tetap bahwa mesin tidak dapat dikaitkan dengan tanggung jawab (dan oleh karena itu juga bukan tidak bertanggung jawab). Namun, manusia dapat bertanggung jawab dan oleh karena itu harus bertanggung jawab atas apa yang mereka lakukan dan putusan ketika mengembangkan atau menggunakan AI.

Meskipun GDPR berlaku untuk sistem AI yang memproses data pribadi, kekosongan regulasi karena tidak mampu mengatasi tantangan unik yang ditimbulkan oleh AI semakin terlihat. Hingga undang-undang diperkenalkan untuk mengatasi tantangan ini, kejelasan dapat diberikan melalui panduan regulasi untuk menerapkan aturan-aturan tertentu secara efektif. Mekanisme regulasi yang efisien harus berfokus pada cara meminimalkan potensi kerugian yang timbul dari keputusan AI, khususnya kerugian yang kemungkinan besar menyebabkan dampak terbesar. Partisipasi pemangku kepentingan yang maksimal dalam struktur tata kelola juga merupakan kunci untuk memastikan bahwa mulai dari konsumen hingga bisnis, peneliti, dan masyarakat sipil, semua pihak terkait harus diajak berkonsultasi tentang bagaimana kerangka regulasi dapat berkembang.

Kesulitan dalam menentukan akuntabilitas telah menyebabkan adanya 'kesenjangan akuntabilitas'. Komisi mengakui bahwa '[b]anyak aktor terlibat' dalam siklus hidup AI, termasuk pengembang, produsen, importir/eksportir, penyebar, dan pengguna. Namun, contoh berikut menyatakan bahwa:

Pengembang AI mungkin berada pada posisi terbaik untuk mengatasi risiko yang timbul dari fase pengembangan, [tetapi] kemampuan mereka untuk mengendalikan risiko selama fase penggunaan mungkin lebih terbatas. Dalam hal ini, penyebar harus tunduk pada kewajiban yang relevan.

Hal ini masih jauh dari berkembang dan tampaknya mengantisipasi bahwa beberapa pengaruh potensial tidak cukup untuk meneruskan kewajiban apa pun bagi pengembang ke tahap penyebaran. Mungkin lebih baik untuk memperlakukan upaya ini sebagai usaha bersama, alih-alih sebagai tongkat estafet (atau kentang panas) yang harus dilimpahkan. Pendekatan ini digunakan untuk isu-isu pertanggungjawaban dalam proposal (di mana pengembang secara teoritis dapat bertanggung jawab atas kerugian yang ditimbulkan dalam contoh yang diberikan). Hal ini, menurutnya, membingungkan karena membayangkan kewajiban satu pihak dengan kewajiban banyak pihak. Hal ini sebagian disebabkan oleh fakta bahwa isu-isu

pertanggungjawaban, sebagaimana dibahas di atas, juga berada dalam aturan Negara Anggota.

Tingkat pengawasan manusia yang diperlukan belum sepenuhnya jelas. Buku Putih mengeksplorasi—tidak secara menyeluruh—tinjauan sebelum penerapan, tinjauan setelah tindakan AI, pemantauan secara menyeluruh, atau pengintegrasian kendala operasional (misalnya penghentian otomatis). Parlemen Uni Eropa telah lebih jelas menyatakan bahwa untuk aplikasi AI berisiko tinggi, pengawasan manusia harus dimungkinkan di semua tahap.

1.5 INVESTASI DAN KOORDINASI

Sebagaimana telah disebutkan, kekurangan dalam menarik investasi dan mengembangkan perusahaan di Uni Eropa telah menyebabkan perlunya tindakan di tingkat Uni Eropa. Angka-angka yang menjadi acuan Komisi menunjukkan investasi sebesar Rp. 32 Triliun dalam AI di Eropa pada tahun 2016, dibandingkan dengan sekitar Rp. 121 Triliun di Amerika Utara dan Rp. 65 Triliun di Asia. Angka-angka ini akan semakin terdampak oleh Brexit, yang pada tahun 2018 tercatat bahwa Inggris menyumbang sepertiga dari seluruh perusahaan AI di Uni Eropa. Hal ini, tentu saja, bukan hanya kerugian yang signifikan, tetapi sekaligus kemunculan instan pesaing yang kuat.

Oleh karena itu, tidak mengherankan bahwa proposal untuk menarik dana sebesar Rp. 200 Triliun untuk AI di Uni Eropa telah disambut baik oleh industri, dengan dukungan khusus untuk investasi dalam keterampilan, inovasi dan penelitian, dan bekerja sama dengan Negara-negara Anggota.

Hal ini mengarah pada sebuah hal yang menarik. Hal ini muncul dalam bentuk keengganan yang tercermin dalam konsultasi mengenai proyek unggulan Komisi, yaitu 'mercusuar'. Komisi menegaskan ada 'kebutuhan' untuk 'pusat penelitian, inovasi, dan keahlian yang akan mengoordinasikan upaya [nasional] ini dan menjadi rujukan dunia untuk keunggulan dalam AI dan yang dapat menarik investasi dan bakat terbaik di bidang tersebut'. Ini disarankan sebagai cara untuk mengatasi 'lanskap pusat kompetensi yang terfragmentasi saat ini yang tidak ada yang mencapai skala yang diperlukan untuk bersaing dengan lembaga-lembaga terkemuka secara global'. Pada tahun 2019 Forbes mencatat bahwa 32 dari 50 perusahaan AI yang paling menjanjikan berpusat di California (mencerminkan pengelompokan yang kuat di sekitar Lembah Silikon). Meskipun pengalaman Tiongkok saat ini lebih tersebar, Tiongkok juga berencana mengembangkan model hub sentral. Keengganan untuk menyatukan bakat dan memfokuskan inovasi di 'mercusuar' UE (tercermin dari 86% yang menyatakan sangat penting untuk mendukung pusat dan jaringan saat ini dengan 'hanya' 64% yang menyatakan skema mercusuar sangat penting) mungkin diharapkan mengingat kepentingan nasional, tetapi itu berarti bahwa pendakian yang lebih organik dan kurang bertenaga dapat terjadi.

Ada juga kekhawatiran yang lebih luas. Seseorang perlu ragu-ragu dalam menarik kesimpulan terlalu cepat dari pandemi virus corona, tetapi adil untuk mengatakan bahwa bahkan di antara negara-negara dan lembaga-lembaga di UE, pengalaman tersebut tidak selalu mendidik mengenai tindakan bersama. Hari terkoordinasi yang direncanakan untuk memulai

vaksinasi COVID massal secara serentak di seluruh Negara Anggota Uni Eropa, dalam apa yang disebut Uni Eropa sebagai 'momen persatuan yang menyentuh', dirusak oleh beberapa negara yang bertindak secara sepihak sebelumnya. Ancaman Komisi untuk menggunakan Pasal 16 protokol Irlandia Utara guna mencegah vaksin tiba di Inggris melalui 'pintu belakang' disambut dengan keraguan baik di Irlandia Utara maupun Republik Irlandia (Negara Anggota Uni Eropa yang frustrasi karena kurangnya konsultasi). Pengawasan Komisi terhadap peluncuran vaksinasi yang lambat juga telah diakui oleh Presiden Komisi von der Leyen dan telah menyebabkan beberapa pihak menyerukan pengunduran diri Presiden. Setiap kekurangan yang dirasakan dalam kepemimpinan Uni Eropa terkait krisis ini dapat memiliki efek limpahan, yang berdampak pada proyek-proyek Uni Eropa yang terkoordinasi.

1.6 CAKUPAN REGULASI DI MASA DEPAN?

Peringatan penting *mungkin* diperlukan terkait implikasi Buku Putih terhadap transparansi, kemudahan penjelasan, dan akuntabilitas perusahaan. Komisi telah menyarankan diferensiasi antarperusahaan berdasarkan risiko. Aplikasi yang dikategorikan sebagai 'berisiko tinggi' diwajibkan untuk mematuhi peraturan yang mengatur AI. Namun, aplikasi yang tidak termasuk dalam kategori ini saat ini tidak akan tercakup dalam peraturan tersebut. Perusahaan-perusahaan tersebut akan memiliki opsi untuk secara sukarela ikut serta dalam skema ini dan kemudian akan menerima 'label mutu' untuk aplikasi AI mereka.

"Risiko tinggi" akan didefinisikan berdasarkan kriteria *kumulatif*. Kriteria pertama menyangkut sifat sektor dan apakah sektor tersebut merupakan sektor yang 'mengingat karakteristik kegiatan yang biasanya dilakukan, risiko signifikan dapat diperkirakan akan terjadi'. Demi kejelasan, diusulkan agar daftar lengkap disediakan dan ditinjau secara berkala. Kriteria kumulatif kedua mensyaratkan bahwa kegiatan tersebut juga merupakan kegiatan di mana penggunaan spesifik AI sedemikian rupa sehingga 'risiko signifikan kemungkinan akan muncul'. Hal ini, disarankan, dapat didasarkan pada dampak yang mungkin terjadi, terutama kematian, cedera, dan kerusakan immaterial, dampak hukum atau yang serupa, dan hasil yang tidak dapat dihindari secara wajar oleh subjek.

Perbedaan antara aplikasi berisiko tinggi (terregulasi) dan aplikasi lainnya (tidak terregulasi) dapat dianggap sebagai dasar yang aman untuk menyusun perencanaan masa depan (atau mungkin sebuah bab dalam kumpulan yang disunting yang membahas perencanaan masa depan). Namun, hanya 43% responden yang mendukung pembatasan regulasi hanya untuk aplikasi berisiko tinggi dalam proposal. Konsultasi publik juga kurang melibatkan publik terkait apakah definisi 'risiko' sudah memuaskan. Oleh karena itu, aspek proposal ini tampaknya tidak memiliki landasan yang kuat.

1.7 JALUR LEGISLATIF

Sejauh ini, bab ini telah membahas tantangan-tantangan spesifik yang disajikan dalam proposal-proposal tersebut. Pada titik ini, akan bermanfaat untuk merenungkan proses yang harus ditempuh agar suatu undang-undang dapat terwujud.

Komisi Eropa-lah yang memegang hak tunggal untuk memprakarsai legislasi di Uni Eropa. Namun, terlepas dari Buku Putih, Parlemen Eropa telah memanfaatkan wewenangannya untuk mengundang Komisi mengusulkan legislasi, dengan pemungutan suara mayoritas untuk mengajukan inisiatif legislatifnya sendiri. Komisi tidak perlu menindaklanjuti hal-hal ini, tetapi harus menanggapi dan memberikan alasan jika memutuskan untuk tidak melakukannya.

Tiga inisiatif tersebut berkaitan dengan pembentukan kerangka etika, penyediaan pertanggungjawaban (dan kewajiban asuransi), dan hak kekayaan intelektual. Proposal mengenai kerangka etika sebagian besar tumpang tindih dengan Buku Putih, sementara proposal pertanggungjawaban dan kekayaan intelektual lebih luas daripada yang terdapat dalam Buku Putih. Pertanggungjawaban, sebagaimana telah dibahas di atas, memiliki perbedaan yang signifikan di antara Negara Anggota dan sulit untuk diselaraskan. Kekayaan intelektual tidak menjadi fokus Buku Putih, dan di sini Parlemen menyarankan bahwa hak kekayaan intelektual hanya boleh diberikan kepada manusia dan bahwa kehati-hatian perlu diambil untuk melindungi inovasi manusia dan prinsip-prinsip etika Uni Eropa.

Penggunaan inisiatif legislatif Parlemen tidak terlalu umum, hanya digunakan 29 kali antara tahun 2009 dan 2019. Kegunaan penerapannya ketika Buku Putih sudah ada dan di mana proses legislatif mengenai masalah tersebut—di mana Parlemen sendiri dapat membuat amandemen—diduga akan segera terjadi agak dipertanyakan. Terutama mengingat bahwa ada tingkat keberhasilan yang relatif rendah dalam mendapatkan persetujuan Komisi atas inisiatif Parlemen tersebut (dengan 7 dari 29 yang disebutkan di atas menghasilkan tindakan Komisi). Siaran pers Parlemen berjudul 'Parlemen memimpin jalan pada set pertama aturan UE untuk Kecerdasan Buatan' berpotensi menyesatkan dan menyajikan basis kekuatan alternatif yang menarik (dan bisa dibilang bermasalah) dalam prosesnya.

Posisi kelembagaan seperti itu, yang terjadi pada Oktober 2020, tidak ada artinya jika dibandingkan dengan pandemi yang sedang berlangsung. Sebagaimana telah ditunjukkan di atas, kekurangan dalam pengelolaan stok vaksinasi dan kontroversi mengenai usulan embargo vaksin yang masuk ke Irlandia Utara—proposal yang kemudian ditolak—telah (pada saat tulisan ini dibuat) menjadi isu yang menonjol, menimbulkan beberapa pertanyaan mengenai kompetensi Komisi. Paling banter, hal ini tidak membantu kemajuan inisiatif Uni Eropa mana pun, dan kenyataannya untuk sebagian besar topik, saat ini pikiran terfokus ke hal lain; AI pun demikian.

Konteks ini, ditambah dengan kompleksitas yang diidentifikasi terkait regulasi AI di Uni Eropa di atas, berarti bahwa regulasi yang akan segera diberlakukan tidaklah realistis. Sebuah kisah peringatan dapat ditemukan dalam Peraturan Perlindungan Data Umum yang agak sebanding, yang membutuhkan waktu sekitar delapan tahun dari konsultasi hingga penerapan.

1.8 PENGADILAN

Prospek periode yang panjang sebelum undang-undang AI khusus disahkan menimbulkan pertanyaan tentang seperti apa bentuk sementara dari 'regulasi' tersebut. Meskipun aturan Uni Eropa yang relevan di atas yang saat ini berlaku untuk AI telah dirujuk

dalam konteksnya, aturan yang tersisa akan berasal dari Negara Anggota itu sendiri. Namun, prospek (saat ini) terdapat 27 perangkat aturan yang berbeda-beda, yang diminimalkan karena peran penting Mahkamah Eropa terkait aturan pasar internal. Kenyataannya, aturan Negara Anggota cenderung menciptakan hambatan terhadap akses pasar barang dan jasa (terutama karena kriteria ini ditafsirkan secara luas), dan aturan tersebut harus dibenarkan secara objektif oleh Negara Anggota atau dihapuskan.

Arah yang diambil Mahkamah terkait isu-isu AI apa pun yang dihadapinya dalam periode interim, menurutnya, juga kemungkinan dipengaruhi oleh posisi yang saat ini diadopsi oleh lembaga-lembaga Uni Eropa (khususnya Komisi). Hubungan simbiosis antara legislatif Uni Eropa dan Mahkamah merupakan hubungan yang mengakar dan mencerminkan realitas di mana harmonisasi negatif (melalui yurisprudensi Mahkamah) sama pentingnya dengan harmonisasi positif (melalui legislasi sekunder). Akibatnya, baik legislatif maupun Mahkamah tampaknya mengejar tujuan yang sama tetapi melalui cara yang berbeda.

Sebagai contoh, hal ini terlihat jelas terkait dengan pergerakan bebas barang dan upaya untuk menghapus hambatan perdagangan non-diskriminatif. Sementara harmonisasi legislasi masih ditunggu, Mahkamah turun tangan dan melalui putusannya di *Cassis de Dijon*, menetapkan prinsip pengakuan bersama—yang pada dasarnya membalikkan keadaan—dan mengizinkan akses pasar tanpa alasan kuat yang sebaliknya.

Sebagaimana dinyatakan di atas, terkait AI, kemungkinan ketentuan pergerakan bebas Uni Eropa akan diterapkan karena langkah-langkah nasional kemungkinan akan menciptakan hambatan perdagangan. Pada titik ini, Mahkamah dapat mengambil kendali tingkat tinggi melalui penggunaan uji proporsionalitas, yang diakui sangat fleksibel. Analisis justifikasi Negara Anggota atas setiap hambatan perdagangan di hadapan Mahkamah berpotensi ketat. Negara Anggota harus menunjukkan kesesuaian, kebutuhan, dan proporsionalitas langkah tersebut. Secara historis, diakui bahwa hal ini terjadi dengan latar belakang di mana Mahkamah mengadopsi pendekatan yang bertujuan untuk interpretasi—pendekatan yang mendukung integrasi dan, akibatnya, penghapusan hambatan perdagangan.

Terlihat pula bahwa intensitas analisis proporsionalitas bervariasi sesuai dengan pokok bahasannya. Misalnya, meningkatnya pentingnya perlindungan lingkungan dan perlindungan konsumen dapat dilacak melalui pelonggaran intensitas peninjauan yurisprudensi. Hal ini pada dasarnya membuat Negara Anggota lebih mudah untuk membenarkan hambatan perdagangan atas dasar ini. Namun, menariknya untuk tujuan kita, perlu dicatat bahwa kebalikannya juga bisa terjadi. Bisa jadi aspek tertentu dari akses pasar atau area perdagangan yang menjanjikan dapat memunculkan analisis proporsionalitas yang lebih intrusif, sehingga membuat pembenaran Negara Anggota menjadi lebih menantang.

Hal ini tampaknya muncul terkait akses ke produk farmasi melalui internet. Dalam kasus *DocMorris*, Pengadilan tidak bersedia menerima larangan menyeluruh atas dasar perlindungan kesehatan masyarakat. Sebaliknya, Pengadilan mengamati bahwa:

Pembelian daring mungkin memiliki keuntungan tertentu, seperti kemampuan untuk memesan dari rumah atau kantor, tanpa perlu keluar rumah, dan memiliki waktu untuk

memikirkan pertanyaan yang akan diajukan kepada apoteker, dan keuntungan ini harus diperhitungkan.

Oleh karena itu, bobot keuntungan kemajuan teknologi akan memengaruhi pembenaran lainnya. Pada akhirnya, dalam kasus ini, tidak ada pembenaran untuk melarang barang-barang non-resep dan barang-barang resep hanya diizinkan untuk dilarang setelah analisis intensif mengenai risiko spesifik bagi individu, dengan argumen yang diajukan oleh pemerintah Jerman, mengenai jaminan sosial dan integritas sistem kesehatan Jerman, ditolak.

Perlu ditegaskan bahwa keputusan ini diambil pada tahun 2003, sebelum belanja daring semakin meluas. Dengan demikian, keputusan ini menunjukkan (i) beberapa pemikiran ke depan dari Pengadilan mengenai potensi teknologi dan (ii) keengganan untuk menyerahkan regulasi kepada Negara Anggota. Ini bukan berarti Pengadilan merupakan garda terdepan tunggal; memang, sudah ada Arahan Uni Eropa yang mengatur penjualan jarak jauh. Namun, Arahan tersebut secara khusus mengizinkan Negara Anggota untuk mengadopsi standar yang lebih ketat guna melindungi konsumen, meskipun Pengadilan mengakui bahwa hak tersebut masih tunduk pada ketentuan pergerakan bebas yang tersisa.

Kasus-kasus ini merupakan indikasi kapasitas Pengadilan untuk memajukan kepentingan Uni Eropa melalui integrasi bagian-bagian pasar yang menguntungkan, baik sebelum maupun setelah undang-undang disahkan. Oleh karena itu, pandangan Pengadilan akan tetap penting untuk waktu yang lama.

Oleh karena itu, meskipun referensi spesifik tentang kecerdasan buatan dalam yurisprudensi sejauh ini bersifat sporadis dan sedikit, patut dicatat bahwa referensi tersebut tampak mendukung.

Tidak mengherankan, referensi tentang AI lebih sering muncul dalam Opini Advokat Jenderal yang memberi nasihat kepada Mahkamah, di mana pendekatan penalaran yang lebih diskursif diadopsi daripada yang diambil oleh Mahkamah itu sendiri. Sejauh ini, hal-hal tersebut menunjukkan kesadaran yang mendalam akan potensi dan pentingnya strategis AI bagi Uni Eropa ke depannya. Bahkan, seorang Advokat Jenderal baru-baru ini menunjukkan, "peluang untuk penyimpanan virtual atau program kecerdasan buatan mau tidak mau mengubah cara profesi dan praktiknya dipahami". Kasus ini menyangkut sektor hukum (yang agak tradisional) dan merupakan pengingat akan jangkauan prospektif AI, dan kepekaan peradilan terhadap hal ini. Advokat Jenderal lainnya menyoroti fakta bahwa definisi istilah 'produk' mungkin perlu ditinjau kembali mengingat perkembangan teknologi modern dan mengutip laporan Komisi mengenai hal ini dengan persetujuan. Sementara itu, Pengadilan Umum enggan membatasi cakupan definisi 'pintar' hanya pada manusia.

Secara praktis, tampaknya pengembangan AI di Uni Eropa akan tetap lebih teratomisasi daripada yang diharapkan Komisi, dengan Negara-negara Anggota dan peserta menilai pusat dan jaringan yang sudah mapan di atas 'lighthouseing'. Tampaknya juga terdapat tekanan yang signifikan untuk meninjau pendekatan risiko yang diusulkan (yang akan membuat banyak aplikasi berada di luar kewenangannya kecuali mereka yang bertanggung jawab berpartisipasi secara sukarela). Revisi ini juga akan membantu mengatasi apa yang masih menjadi kekhawatiran mendasar yang signifikan terkait AI di Uni Eropa secara umum.

Lebih mendasar lagi, seiring dengan semakin banyaknya sistem AI yang terintegrasi ke dalam kehidupan sehari-hari, diprediksi bahwa AI akan memiliki efek transformatif global terhadap struktur ekonomi dan sosial, serupa dengan efek teknologi umum lainnya, seperti mesin uap, kereta api, listrik, elektronik, dan internet.

Meskipun mudah untuk terjebak dalam aspek inovatif AI, AI masih memiliki jalan panjang sebelum dapat dianggap aman bagi semua pihak yang terlibat. Penggunaan data yang luas merupakan pedang bermata dua, yang memungkinkan mesin untuk membuat keputusan lebih cepat tetapi juga mengungkap beberapa kerentanan yang berpotensi menimbulkan kompromi etis. Menganalisis berbagai tantangan etika yang ditimbulkan oleh AI, jelas bahwa transparansi, keterjelasan, dan akuntabilitas saling terkait dan terikat oleh tema umum kepercayaan. Komisi dalam Buku Putihnya menyebut hal ini sebagai 'ekosistem kepercayaan', yang seharusnya memberikan warga negara kepercayaan untuk menggunakan sistem AI dan perusahaan serta badan publik kepastian hukum untuk berinovasi menggunakan AI.

Sangat penting bagi regulator dan bisnis untuk berbagi visi membangun kepercayaan ini dengan menerapkan AI dengan kerangka etika yang baik. Konsumen juga penting, dan agar AI berhasil, individu yang terdampak harus yakin bahwa hak-hak fundamental mereka tidak dikompromikan. Menempatkan kepercayaan sebagai unsur yang tak tergantikan, bisnis perlu mengidentifikasi manfaat spesifik yang dapat dibawa oleh teknologi AI dan menimbangnya dengan biaya penerapannya secara bertanggung jawab dan etis.

Dengan banyaknya bisnis yang menerapkan atau mengeksplorasi opsi AI, mungkin keliru dipercaya bahwa keberhasilan AI adalah masalah keuntungan finansial. Namun, yang jauh lebih penting adalah bagaimana AI terhubung langsung dengan manusia yang terdampak. Menempatkan kesejahteraan manusia sebagai inti pembangunan menetapkan tujuan yang realistis serta cara konkret untuk mengukur dampak AI. Komisi berpandangan bahwa agensi dan pengawasan manusia adalah elemen kunci dari sistem AI yang dapat dipercaya. AI harus memberdayakan manusia, memungkinkan mereka untuk membuat keputusan yang terinformasi sambil melindungi hak-hak fundamental. Unsur 'manusia' perlu hadir dalam semua aspek AI, seperti keterlibatan manusia, keterlibatan manusia dalam proses, dan peran manusia dalam komando.

Komisi berpandangan bahwa keterlibatan manusia adalah kunci untuk memastikan otonomi manusia tidak terabaikan, dan tujuan AI yang tepercaya dan etis tetap utuh. Pengawasan manusia dipandang penting dalam beberapa aspek penerapannya, mulai dari pemantauan dan perancangan hingga keluaran sistem AI. Namun, orang mungkin bertanya-tanya, jika tinjauan manusia diperlukan untuk mencapai hasil AI yang efektif, apakah kita benar-benar mencapai dampak yang dijanjikan AI?

Menyelesaikan masalah ini akan membutuhkan waktu. Sementara itu, perusahaan dapat yakin bahwa Pengadilan tidak akan membiarkan munculnya peraturan nasional yang berbeda. Namun, selain keengganan untuk mengizinkan Negara Anggota mengatur sektor ini, Pengadilan akan memiliki — dan seharusnya memiliki — lebih sedikit hal untuk dikatakan dan kemungkinan akan dipandu oleh proses legislatif. Hal ini disambut baik karena, saat ini, hal terakhir yang dibutuhkan AI di Uni Eropa adalah suara yang berbeda.

BAB 2

DAMPAK PENGENALAN WAJAH AI TERHADAP PRIVASI

2.1 KONSEP DAN FUNGSI FRT

Sebagian besar definisi 'teknologi pengenalan wajah' (FRT) dipengaruhi oleh berbagai fungsi yang dijalankan oleh teknologi tersebut. Terdapat tiga fungsi FRT yang berbeda. Pertama adalah 'identifikasi', terkadang disebut sebagai 'perbandingan satu-ke-banyak'. 'Identifikasi' berarti templat citra wajah seseorang dibandingkan dengan banyak templat lain yang tersimpan dalam basis data. FRT kemudian menghasilkan skor untuk setiap perbandingan, yang menunjukkan kemungkinan bahwa dua citra merujuk pada orang yang sama. Akibatnya, istilah 'pengenalan wajah otomatis' (AFR) terkadang digunakan untuk merujuk pada proses ini. Fungsi kedua dikenal sebagai 'verifikasi' atau 'otentikasi', juga disebut sebagai 'perbandingan satu-ke-satu'. Aplikasi ini membandingkan dua templat biometrik untuk menentukan kemungkinan bahwa dua gambar menunjukkan orang yang sama. Jika kemungkinannya di atas ambang batas tertentu, identitas diverifikasi. Ini pada dasarnya adalah operasi identifikasi. Namun, alih-alih dilakukan terhadap semua templat wajah yang disimpan dalam basis data, itu hanya terhadap templat unik (misalnya, ketika seseorang membuka kunci ponsel mereka menggunakan pengenalan wajah). Akhirnya, FRT juga dapat digunakan untuk apa yang disebut 'kategorisasi', yaitu analisis wajah. Dalam aplikasi ini, teknologi ini tidak digunakan untuk mengidentifikasi atau mencocokkan individu, tetapi untuk mendapatkan karakteristik tertentu dari individu, yang tidak selalu memungkinkan untuk identifikasi. Karakteristik tersebut mungkin orientasi seksual, jenis kelamin, usia, kesehatan mental, potensi untuk melakukan kejahatan, dan banyak lagi.

2.2 BAGAIMANA AI MEMENGARUHI FRT?

Sebelumnya, FRT terbatas pada fungsi identifikasi dan verifikasi (fungsi pertama dan kedua yang telah disebutkan). Lebih lanjut, citra wajah yang dimasukkan ke dalam sistem harus memenuhi persyaratan yang sangat spesifik. Misalnya, foto harus digambarkan dalam 'gaya ID', dengan subjek menghadap ke depan dan menatap langsung ke kamera. Namun, teknologi berkembang, terutama dalam hal resolusi. Akhirnya, AI ditambahkan, yang meningkatkan akurasi identifikasi dan verifikasi. Misalnya, kini AI dapat mengidentifikasi wajah yang bergerak, saat tidak menatap langsung ke kamera, dan bahkan individu yang mengenakan aksesoris seperti kacamata atau topi. Aspek paling kontroversial yang dibawa oleh AI mungkin berkaitan dengan kategorisasi. Berikut ini adalah daftar yang tidak lengkap dari beberapa contoh dampak AI pada FRT:

- Wang dan Kosinski mengklaim telah mencapai sistem pengenalan wajah menggunakan jaringan saraf untuk mengekstrak orientasi seksual dari citra wajah.
- Ada juga beberapa proyek penelitian yang bekerja pada deteksi depresi dari citra wajah, seperti studi yang dilakukan oleh Zhou, Jin, Shang, dan Guo. Mereka

menggunakan jaringan saraf konvolusional dalam untuk menghasilkan skor berdasarkan, apa yang mereka sebut, 'peta aktivasi depresi yang dihasilkan (DAM)'.

- FRT juga digunakan dalam proyek iBorderCtrl dalam wawancara pribadi dengan para pelancong. Tujuan proyek ini adalah untuk mempercepat prosedur penyeberangan perbatasan. Untuk mencapai tujuan ini, FRT digunakan untuk mendeteksi apakah orang berbohong untuk mencegah penyeberangan ilegal. Pemilihan FRT didasarkan pada kriteria skalabilitas, pengurangan beban kerja, dan tujuan untuk menghindari kesalahan yang dibuat oleh agen manusia. Hal ini juga didukung oleh fakta bahwa FRT merupakan teknologi non-intrusif karena tidak terlihat oleh individu yang dikenainya.

Contoh-contoh ini memberikan gambaran yang lebih baik tentang ranah yang telah dibuka oleh aplikasi AI untuk fungsi kategorisasi FRT.

2.3 APAKAH CITRA WAJAH MERUPAKAN DATA PRIBADI?

Citra wajah adalah data masukan yang diterima sistem FRT dan dengan demikian melampirkan nama dan nama keluarga (dalam hal identifikasi dan verifikasi) atau ciri lain (dalam hal kategorisasi). Karena FRT diisi dengan citra wajah, penting untuk membedakan apakah citra-citra ini merupakan data pribadi dan, akibatnya, apakah rezim hukum tertentu, seperti Peraturan Perlindungan Data Umum (GDPR) Uni Eropa, berlaku. Peraturan Perlindungan Data Umum dianggap sebagai salah satu undang-undang perlindungan data terpenting di dunia dan, meskipun cakupannya regional (UE), peraturan perlindungan data di seluruh dunia menjadikan GDPR sebagai panutan. Istilah 'data pribadi' didefinisikan dalam Pasal 4.1 GDPR sebagai:

setiap informasi yang berkaitan dengan orang perseorangan yang teridentifikasi atau dapat diidentifikasi ('subjek data'); orang perseorangan yang dapat diidentifikasi adalah orang yang dapat diidentifikasi, secara langsung maupun tidak langsung, khususnya dengan mengacu pada pengenal seperti nama, nomor identifikasi, data lokasi, pengenal daring, atau satu atau lebih faktor khusus yang berkaitan dengan identitas fisik, fisiologis, genetik, mental, ekonomi, budaya, atau sosial orang perseorangan tersebut.

Pada awalnya, citra wajah dapat dianggap sebagai data pribadi, selama citra tersebut memungkinkan identifikasi seseorang. Hanya dengan citra wajah (yang relatif mudah diperoleh, misalnya, dari jejaring sosial), sejumlah besar data pribadi (yang sangat) dapat diperoleh dengan menggunakan FRT. Namun, citra wajah juga dapat diproses dan/atau ditambahkan *noise* sehingga mustahil untuk menyimpulkan identitas orang tertentu dari citra tersebut. Jika demikian, citra wajah tidak akan dianggap sebagai data pribadi, melainkan data non-pribadi.

Aspek penting lainnya, yang seringkali luput dari perhatian, yang perlu dipertimbangkan ketika menganalisis apakah citra wajah merupakan data pribadi atau bukan adalah potensi pembalikannya untuk memungkinkan identifikasi. Jika kita menganggap citra wajah yang digunakan untuk mengidentifikasi subjek sebagai data pribadi, kita berada dalam ranah perlindungan data dan, oleh karena itu, GDPR berlaku. Dalam GDPR, tidak semua data

pribadi dianggap berada pada 'tingkat' yang sama dan terdapat kategori khusus yang disebut 'data sensitif'.

Data ini berpotensi mengungkapkan 'asal ras atau etnis, opini politik, keyakinan agama atau filosofis, atau keanggotaan serikat pekerja, [...] data genetik, data biometrik, [...] data mengenai kesehatan atau [...] kehidupan seks atau orientasi seksual seseorang'. Karena 'sifatnya yang sensitif', pemrosesan data ini tidak diizinkan berdasarkan GDPR, kecuali dalam kasus-kasus tertentu yang diatur dalam undang-undang. Data biometrik (gambar wajah yang digunakan untuk mengidentifikasi seseorang) secara otomatis dianggap sebagai data sensitif.

Akibatnya, gambar wajah dapat dianggap sebagai 'data pribadi' ketika diperoleh untuk melakukan identifikasi dan verifikasi menggunakan FRT. Selain itu, banyak aplikasi FRT saat ini yang menjalankan fungsi kategorisasi, sebagai akibat dari kinerjanya, akan mengungkapkan data yang sangat pribadi, seperti orientasi seksual, kondisi kesehatan mental, dan banyak lagi. Ini juga dapat dianggap sebagai data pribadi. Penggunaan FRT untuk melakukan kategorisasi telah menerima banyak reaksi keras dalam studi ini. Para akademisi percaya bahwa bidang studi yang dikenal sebagai 'analisis sentimen' (juga disebut sebagai 'AI fisiognomis') ini tidak memiliki dasar ilmiah dan, oleh karena itu, tidak valid.

Meskipun FRT (dan AI) tidak disebutkan secara eksplisit dalam GDPR, banyak ketentuan yang berlaku untuk teknologi tersebut. Sartor berpendapat bahwa pengenalan istilah-istilah terkait Internet dalam GDPR merupakan reaksi terhadap ketiadaan istilah-istilah ini dalam Arahan Perlindungan Data sebelumnya karena konteks historis dan sosialnya. Oleh karena itu, tampaknya GDPR telah gagal memasukkan istilah-istilah terkait AI dengan cara yang sama. Hal ini mungkin disebabkan oleh sifat perkembangan AI terkini yang relatif baru dan inovatif, yang tidak diprediksi pada saat GDPR disusun. Hal ini mungkin juga berkaitan dengan fakta bahwa GDPR tidak berorientasi pada teknologi apa pun secara eksklusif, melainkan bertujuan untuk mengatasi semua ancaman terhadap hak atas perlindungan data terlepas dari asal usulnya. Secara umum, para akademisi berpendapat bahwa GDPR gagal memberikan panduan tentang cara menerapkan teknologi yang ditingkatkan AI dengan tetap menghormati isinya. Kritik utama difokuskan pada isinya yang luas dan seringkali samar.

2.4 AI DAN PRIVASI: PENGENALAN WAJAH DI INDONESIA

Penggunaan teknologi pengenalan wajah atau Face Recognition Technology (FRT) di Indonesia telah membuka peluang besar dalam peningkatan keamanan publik dan penegakan hukum, namun juga menimbulkan tantangan serius terkait privasi dan perlindungan data pribadi. Meskipun Indonesia telah memiliki dasar hukum melalui Undang-Undang Nomor 27 Tahun 2022 tentang Perlindungan Data Pribadi (UU PDP), regulasi ini masih dianggap belum cukup spesifik dalam mengatur penggunaan teknologi FRT secara rinci. Hal ini terutama terlihat pada praktik pemasangan kamera pengawas yang dilengkapi FRT di ruang publik maupun institusi tertentu, di mana individu yang terekam sering kali tidak memberikan persetujuan eksplisit untuk pemrosesan data biometriknya. Selain itu, penggunaan FRT dalam konteks penegakan hukum atau keamanan publik membawa risiko penyalahgunaan data. Teknologi ini berpotensi digunakan untuk identifikasi yang keliru, pelacakan individu tanpa dasar hukum, atau pemantauan perilaku warga secara masif, yang dapat menimbulkan

pelanggaran hak privasi dan kebebasan sipil. Kurangnya transparansi mengenai bagaimana data wajah dikumpulkan, disimpan, dan digunakan memperburuk potensi risiko tersebut, terutama bila data jatuh ke tangan pihak yang tidak bertanggung jawab atau diproses tanpa prosedur keamanan yang memadai.

Oleh karena itu, diperlukan regulasi khusus untuk FRT di Indonesia. Regulasi tersebut perlu mencakup aspek persetujuan individu, keamanan data biometrik, pembatasan penggunaan oleh pihak berwenang, serta mekanisme pengawasan dan akuntabilitas yang jelas. Tujuannya adalah memastikan bahwa teknologi FRT dapat dimanfaatkan untuk tujuan keamanan dan penegakan hukum secara etis, tanpa mengorbankan hak privasi dan perlindungan data pribadi warga negara. Pendekatan regulatif yang matang akan membantu mencegah potensi penyalahgunaan, membangun kepercayaan publik terhadap penggunaan AI, dan mengintegrasikan teknologi canggih ini ke dalam sistem hukum dan keamanan Indonesia secara aman dan bertanggung jawab.

2.5 TANTANGAN HUKUM TERKAIT FRT

FRT menimbulkan beberapa tantangan hukum: mulai dari keadilan algoritmik, transparansi, dan akuntabilitas hingga hak atas privasi dan perlindungan data. Karena teknologi ini telah mengalami peningkatan eksponensial dalam penerapannya selama 5 tahun terakhir, tantangan-tantangan ini menjadi semakin nyata. Darurat kesehatan COVID-19, misalnya, memerlukan penerapan FRT secara luas. Sebagai teknologi nirsentuh dan tidak mengganggu, FRT sangat cocok untuk memantau kepatuhan terhadap tindakan karantina atau karantina wilayah, atau kewajiban untuk mengenakan masker. Oleh karena itu, semakin penting untuk mengatasi tantangan hukum yang terkait dengan teknologi ini guna menghindari pelanggaran hak asasi manusia dan untuk menegakkan penerapan FRT yang sah.

Dari sudut pandang keadilan algoritmik, FRT mungkin bias terhadap karakteristik tertentu, termasuk ras, gender, dan lainnya. Sudah diketahui umum bahwa algoritma pengenalan wajah masih kesulitan mengidentifikasi orang kulit berwarna. Lebih lanjut, literatur telah menyoroti kerugian yang ditimbulkan oleh penggunaan FRT untuk identifikasi/verifikasi bagi orang transgender dan non-biner. Hal ini disebabkan oleh kumpulan data pelatihan yang terdiri dari citra wajah yang sebagian besar dimiliki oleh orang kulit putih, terutama pria. Selain itu, FRT masih belum memiliki fondasi yang kuat untuk mengenali wajah yang telah diubah karena kecelakaan atau kelumpuhan atau individu yang telah menjalani prosedur bedah wajah. Jika kumpulan data pelatihan merupakan cerminan masyarakat tempat kita tinggal dan masyarakat tersebut dibangun di atas ketidaksetaraan, kinerja FRT akan melanggengkan ketidaksetaraan tersebut.

Perlu juga dipertimbangkan bahwa FRT, bahkan dalam 'tampilan fisiknya', telah dikembangkan dengan mengambil model subjek 'rata-rata' sebagai referensi. Oleh karena itu, representasi dan penelitian tidak ada tentang bagaimana orang-orang dengan disabilitas atau dengan perbedaan kraniofasial berhubungan dengan FRT. Bidang ini dikenal sebagai 'politik artefak' dan hampir tidak mempelajari FRT, tetapi dapat membantu membuatnya lebih adil. Beberapa berpendapat bahwa FRT sering dirancang untuk mencakup berbagai kepentingan dan mengabaikan yang lain.

Sebagai contoh, Introna dan Wood menyebutkan bahwa mesin ATM bank dirancang gagal untuk mempertimbangkan seseorang di kursi roda atau tidak dapat memasukkan kode pin karena disabilitas. Lebih lanjut, trade-off antara keadilan dan akurasi telah menjadi salah satu hambatan utama untuk perbaikan dalam pengertian ini. Sementara ilmuwan komputer mengklaim bahwa tujuan mereka adalah akurasi dalam sistem FRT, tidak mempertanyakan ketidaksetaraan sosial yang disebutkan dalam paragraf sebelumnya, pengacara lebih fokus pada FRT yang mencerminkan keragaman dan keinginan untuk masyarakat yang lebih adil.

Pertanyaan tentang apakah, dan bagaimana, akuntabilitas dapat dituntut dari algoritma dibahas secara luas dalam literatur. Siapa yang harus dimintai pertanggungjawaban? Berdasarkan kriteria apa, dan kepada siapa? Pengawas Perlindungan Data Eropa telah menyatakan keprihatinannya tentang pertanyaan-pertanyaan ini dan literatur sedang menjajaki kemungkinan pembentukan otoritas, pengawas, untuk FRT. Sejalan dengan hal ini, baru-baru ini terdapat beberapa masukan menarik mengenai penerapan FDA untuk FRT. Klaim utamanya adalah bahwa '[m]enangani trade-off antara risiko dan manfaat FRT yang kompleks memerlukan pembentukan kantor federal baru.' Para penulis sampai pada kesimpulan ini dengan menggambar analogi dengan struktur regulasi untuk industri kompleks lainnya yang telah berhasil ditangani oleh lembaga federal.

Contoh regulasi yang menarik mengenai akuntabilitas adalah RUU Negara Bagian Washington, yang dikhususkan untuk FRT. Regulasi ini hampir seluruhnya berfokus pada apa yang disebut 'laporan akuntabilitas'. Laporan akuntabilitas ini harus diajukan oleh setiap lembaga negara bagian atau lokal yang mengembangkan, mengadakan, atau menggunakan FRT dan menjelaskan alasan penggunaan teknologi tersebut. Laporan ini akan berisi ringkasan tujuan penggunaan layanan, informasi tingkat kecocokan palsu, ketentuan keamanan data, protokol pemantauan, dan saluran umpan balik. Lebih lanjut, setiap 'laporan akuntabilitas' harus melalui proses persetujuan demokratis yang akan direvisi setiap dua tahun. Mengenai transparansi, laporan tersebut harus tersedia di situs web lembaga masing-masing dan diteruskan ke otoritas legislatif yang berwenang untuk dipublikasikan di situs webnya. Sehubungan dengan keadilan, penyedia layanan diharuskan menyediakan antarmuka pemrograman aplikasi untuk memungkinkan pemantauan independen terhadap konsistensi dan disparitas keluaran yang tidak merata di seluruh subpopulasi yang berbeda. Jika temuan penelitian independen mengungkapkan disparitas kinerja yang tidak merata di seluruh subpopulasi, penyedia harus merumuskan dan melaksanakan strategi untuk menyelesaikan perbedaan kinerja yang dilaporkan dalam waktu 90 hari setelah menerima hasilnya. Selain itu, metode dan data yang digunakan dalam penelitian independen harus diungkapkan kepada penyedia dengan cara yang memungkinkan replikasi lengkap untuk mengurangi perbedaan yang dilaporkan. Selain itu, lembaga yang menggunakan layanan tersebut untuk membuat keputusan yang memiliki konsekuensi hukum atau yang sama pentingnya bagi orang-orang harus menilai layanan di lingkungan operasi sebelum mengimplementasikannya, untuk mencapai hasil dengan kualitas tertinggi dengan mengikuti instruksi pengembang layanan. RUU tersebut juga mengatur kewajiban setiap lembaga yang menggunakan FRT untuk mendidik semua orang yang menjalankan teknologi tersebut atau mengelola data pribadi yang

dikumpulkannya secara berkala. Terakhir, RUU tersebut mempertimbangkan pembatasan tertentu seperti larangan menggunakan FRT untuk melakukan deteksi waktu nyata atau hampir waktu nyata, atau memulai pemantauan berkelanjutan kecuali surat perintah dikeluarkan, kondisi mendesak muncul, atau perintah pengadilan diperoleh untuk tujuan tunggal menemukan atau mengenali individu yang hilang atau mengidentifikasi orang yang telah meninggal.

2.6 PRIVASI DAN PERLINDUNGAN DATA

Sebagaimana dijelaskan sebelumnya, citra wajah, makanan untuk FRT, dapat dianggap sebagai data pribadi. Oleh karena itu, terdapat beragam isu yang ditimbulkan FRT terkait hak privasi dan perlindungan data. Dimulai dengan persetujuan, terdapat banyak skenario penerapan di mana individu dapat menjadi sasaran FRT tanpa persetujuan mereka. Misalnya, FRT telah diterapkan di beberapa supermarket di seluruh Eropa. Supermarket tersebut mengklaim bahwa tujuan mereka dalam menerapkan teknologi tersebut adalah untuk mencegah masuknya orang-orang dengan perintah penahanan preventif atau aktif ke tempat tersebut atau pekerjaannya. Namun, penggunaan ini dianggap tidak proporsional oleh beberapa Otoritas Perlindungan Data (DPA) dan akibatnya dilarang. Pasal 9.1 GDPR menyatakan bahwa:

[p]roses data pribadi yang mengungkapkan asal ras atau etnis, opini politik, keyakinan agama atau filosofis, atau keanggotaan serikat pekerja, dan pemrosesan data genetik, data biometrik untuk tujuan mengidentifikasi seseorang secara unik, data mengenai kesehatan atau data mengenai kehidupan seks atau orientasi seksual seseorang dilarang.

Namun, paragraf 2(a) dari ketentuan yang sama menunjukkan bahwa '[p]aragraf 1 tidak berlaku jika salah satu dari berikut ini berlaku: [...] subjek data telah memberikan persetujuan eksplisit untuk pemrosesan data pribadi tersebut untuk satu atau lebih tujuan tertentu.'

Lebih lanjut, Pasal 9.2(e) GDPR menetapkan kemungkinan melakukan 'pemrosesan yang terkait dengan data pribadi, yang secara nyata dipublikasikan oleh subjek data'. Hal ini dapat terjadi pada gambar yang dibagikan secara sukarela di profil publik pada layanan jejaring sosial (SNS). Penggunaan gambar wajah dari SNS ini telah dilaporkan dilakukan oleh Clearview AI. Perusahaan Amerika ini menjual layanan pengenalan wajah yang klien utamanya adalah lembaga penegak hukum di AS. Perusahaan tersebut mengklaim program mereka membantu mengidentifikasi pelaku dan korban kejahatan serta melacak ratusan penjahat yang berkeliaran, termasuk pedofil, teroris, dan pedagang seks. Program ini juga digunakan untuk membantu membebaskan orang yang tidak bersalah dan mengidentifikasi korban kejahatan, termasuk pelecehan seksual anak dan penipuan keuangan.

Mereka menuduh telah membangun basis data gambar wajah yang mengesankan hanya dengan menggunakan informasi publik dari web terbuka. Sejak perusahaan mulai mendapatkan popularitas tertentu, perusahaan telah menghadapi banyak tentangan karena kurangnya transparansi mengenai asal gambar wajah yang digunakan dalam basis data (pelatihannya) serta karena potensinya untuk menjadi alat penyalahgunaan wewenang polisi.

Clearview AI juga baru-baru ini menjadi perhatian dari lembaga penegak hukum di seluruh Uni Eropa. Aktivitas Clearview telah menjadi objek dari beberapa pertanyaan parlemen tentang apakah: (a) layanan mereka digunakan oleh siapa pun di Uni Eropa; (b) mereka menyimpan data warga negara Uni Eropa dan (dalam hal ini) bagaimana data tersebut diproses; dan (c) mereka konsisten dengan peraturan perlindungan data UE dan perjanjian bilateral UE-AS tentang privasi. Namun, beberapa anggota Parlemen Eropa telah menunjukkan ketidaksetujuan mereka dengan tanggapan ambigu Komisi Eropa. Praktik Clearview AI mendorong Dewan Perlindungan Data Eropa (EDPB) untuk mengeluarkan surat tentang masalah ini. Mereka menyimpulkan bahwa 'penggunaan layanan seperti Clearview AI oleh otoritas penegak hukum di Uni Eropa, seperti yang terjadi, kemungkinan tidak konsisten dengan rezim perlindungan data UE'. Sebagai konsekuensinya, pada Januari 2021 DPA Hamburg menganggap profil biometrik Clearview AI milik orang Eropa ilegal. Pengadilan juga memerintahkan perusahaan untuk menghapus nilai hash matematika yang mewakili profil biometrik Matthias Marx, seorang warga Hamburg dan anggota Chaos Computer Club yang profil biometriknya telah ditambahkan ke basis data Clearview yang dapat dicari tanpa sepengetahuannya. Akhirnya, pada Februari 2021, DPA Swedia mendenda Otoritas Kepolisian Swedia setelah menemukan bahwa mereka telah menggunakan Clearview AI untuk mengidentifikasi individu. DPA menyimpulkan bahwa:

Kepolisian telah gagal menerapkan langkah-langkah organisasi yang memadai untuk memastikan dan mampu menunjukkan bahwa pemrosesan data pribadi, dalam kasus ini, telah dilakukan sesuai dengan Undang-Undang Data Kriminal. Saat menggunakan Clearview AI, Kepolisian telah memproses data biometrik secara tidak sah untuk pengenalan wajah serta gagal melakukan penilaian dampak perlindungan data yang diperlukan dalam kasus pemrosesan ini.

Dalam kasus kategorisasi wajah, persetujuan mungkin tidak bebas dan berdasarkan informasi karena penggunaan dan fungsi teknologi tidak akan dijelaskan secara jelas karena sifatnya yang inovatif dan kesulitan intrinsik untuk memahami teknologi yang didukung AI. Oleh karena itu, pemrosesan tersebut dapat dianggap melanggar hukum.

Lebih lanjut, Pasal 7 GDPR menetapkan syarat-syarat yang diperlukan untuk persetujuan. Persetujuan harus diberikan secara bebas, spesifik, berdasarkan informasi, dan tidak ambigu. Demikian pula, Pasal 15 Undang-Undang Privasi Informasi Biometrik (BIPA) Illinois mewajibkan persetujuan berdasarkan informasi untuk pemrosesan data biometrik. Berbagai fungsi yang mungkin dilakukan FRT harus dipertimbangkan dalam hal ini. Misalnya, mungkin terdapat beban yang berbeda bagi subjek data untuk menyetujui agar tunduk pada FRT untuk check-in di bandara, alih-alih salah satu aplikasi yang disebutkan di bagian kedua bab ini. Beberapa penulis sangat mengaitkan kebutuhan untuk mengklarifikasi ketentuan persetujuan dengan prinsip pembatasan tujuan. Mengingat bahwa FRT mutakhir melibatkan AI, prinsip pembatasan tujuan menjadi semakin penting. Akibatnya, pemroses/pengendali data harus menentukan data yang dikumpulkan (citra wajah, templat, nama, dll.), fungsi yang dilakukan oleh FRT (identifikasi/verifikasi atau kategorisasi), apakah fungsi tersembunyi untuk

data yang dikumpulkan dipertimbangkan (seperti membangun basis data citra wajah), dan kebijakan penyimpanan data. Dalam hukum AS, BIPA menyatakan dalam pasal 15(b) bahwa 'tidak ada entitas swasta yang boleh mengumpulkan, menyimpan, atau menggunakan informasi biometrik tanpa terlebih dahulu memberikan pemberitahuan kepada, dan memperoleh rilis tertulis atau persetujuan dari, subjek'.

Lebih lanjut, karena sifat inovatif teknologi ini, dan rendahnya tingkat kepercayaan yang dimilikinya, orang-orang tidak memiliki pengetahuan dan kekuatan yang cukup untuk memahami dampak sebenarnya dari apa yang mereka setuju. Sebagai 'tindakan pencegahan' yang memungkinkan, Pasal 13 dan 14 GDPR menetapkan bahwa subjek data berhak untuk diberitahu tentang pengumpulan dan penggunaan data pribadi mereka. Hal ini memungkinkan subjek data untuk menggunakan hak mereka ketika persetujuan belum diberikan. Lebih lanjut, Pasal 15 GDPR menetapkan bahwa:

Subjek data berhak untuk mendapatkan konfirmasi dari pengendali mengenai apakah data pribadi yang berkaitan dengan dirinya sedang diproses atau tidak, dan, jika demikian, akses ke data pribadi dan informasi berikut: (a) tujuan pemrosesan; (b) kategori data pribadi yang bersangkutan; (c) penerima atau kategori penerima yang kepadanya data pribadi telah atau akan diungkapkan, khususnya penerima di negara ketiga atau organisasi internasional; (d) jika memungkinkan, periode yang direncanakan untuk penyimpanan data pribadi, atau, jika tidak memungkinkan, kriteria yang digunakan untuk menentukan periode tersebut; (e) adanya hak untuk meminta pengawas melakukan perbaikan atau penghapusan data pribadi atau pembatasan pemrosesan data pribadi mengenai subjek data atau untuk menolak pemrosesan tersebut; (f) hak untuk mengajukan keluhan kepada otoritas pengawas; (g) jika data pribadi tidak dikumpulkan dari subjek data, informasi apa pun yang tersedia mengenai sumbernya; (h) adanya pengambilan keputusan otomatis, termasuk pembuatan profil, sebagaimana dimaksud dalam Pasal 22(1) dan (4) dan, setidaknya dalam kasus tersebut, informasi yang bermakna tentang logika yang terlibat, serta signifikansi dan konsekuensi yang diperkirakan dari pemrosesan tersebut bagi subjek data.

Lebih lanjut, ketidakcocokan yang tampak antara Big Data (basis data FRT terdiri dari banyak gambar wajah dari berbagai latar belakang) dan Pasal 9 GDPR mungkin terlihat. Penggunaan Big Data untuk melatih sistem AI dan memungkinkannya membuat inferensi mungkin bertentangan dengan keabsahan pemrosesan karena beberapa hasil pelatihan tidak dapat diantisipasi (seperti adanya bias algoritmik). Oleh karena itu, prinsip-prinsip GDPR dari Pasal 5, khususnya batasan tujuan dan minimisasi data, harus diterapkan dengan cara yang tidak berbenturan dengan kinerja AI. Minimalisasi data dan batasan penyimpanan juga berlaku untuk FRT karena menimbulkan pertanyaan berikut: Di mana citra wajah dan templat disimpan? Siapa yang memiliki akses ke sana? Mungkinkah untuk 'bertukar' atau bahkan menjual basis data citra wajah atau templat?

Terlebih lagi, saat ini industri harus berurusan dengan meluasnya penggunaan masker wajah karena COVID-19. Hal ini kemungkinan akan berlanjut untuk beberapa waktu setelah pandemi. Upaya ini menimbulkan pertanyaan apakah gambar wajah lengkap tidak lagi

diperlukan. Beberapa perusahaan mengklaim mereka telah mencapai tujuan ini. Operasi FRT, seperti identifikasi atau verifikasi dengan gambar wajah, akan dapat dilakukan, di mana orang tersebut mengenakan masker wajah, karena seluruh wajah tidak akan diperlukan untuk pengenalan wajah. Sejalan dengan prinsip minimisasi data, semakin sedikit fitur wajah yang dibutuhkan, semakin baik untuk perlindungan data. Namun, hal ini mungkin menimbulkan kekurangan terkait akurasi sistem FRT, yang meningkatkan jumlah identifikasi positif palsu.

Serangkaian masalah penting lainnya pada persimpangan antara FRT dan privasi/perlindungan data adalah pembuatan profil. Pasal 4.4 GDPR mendefinisikan pembuatan profil sebagai:

segala bentuk pemrosesan data pribadi secara otomatis yang terdiri dari penggunaan data pribadi untuk mengevaluasi aspek-aspek pribadi tertentu yang berkaitan dengan seseorang, khususnya untuk menganalisis atau memprediksi aspek-aspek yang berkaitan dengan kinerja orang tersebut di tempat kerja, situasi ekonomi, kesehatan, preferensi pribadi, minat, keandalan, perilaku, lokasi, atau pergerakannya.

Banyak sistem FRT termasuk dalam definisi ini, seperti sistem yang memantau preferensi pemasaran. Lebih lanjut, keadaan darurat kesehatan COVID-19 telah menyebabkan peningkatan penerapan FRT untuk keperluan pembuatan profil seperti pemantauan suhu (pemindai termal FRT), tetapi juga kinerja dan perhatian karyawan dan siswa (termasuk di rumah).

Hukum dan teknologi tidak berdiri sendiri. Di dunia saat ini, dunia yang semakin bergerak menuju 'tata kelola digital', batas antara kedua disiplin ilmu ini semakin kabur. FRT dengan sempurna menggambarkan teka-teki ini. FRT adalah alat yang memungkinkan hasil yang semakin bermanfaat, menarik, dan tak terduga di berbagai bidang, dan berada dalam bahaya karena potensi ancaman yang ditimbulkannya terhadap hak privasi dan khususnya perlindungan data, serta dari sudut pandang keadilan, transparansi, dan akuntabilitas algoritmik. Teknologi ini menghadirkan tantangan bagi model privasi/perlindungan data yang mungkin tidak lagi relevan dengan skenario saat ini, di mana teknologi yang didukung AI semakin menguat. Satu-satunya cara untuk mengatasi tantangan ini adalah dengan melibatkan para profesional yang sama yang bertanggung jawab untuk merancang dan menerapkan teknologi ini dalam pekerjaan regulasi, dengan menempatkan alat yang sama yang menjadikan teknologi ini inovatif untuk layanan perlindungan hak-hak subjek.

Salah satu solusi yang mungkin untuk dilema persetujuan adalah menetapkan standar kepatuhan untuk persetujuan. Standar ini dapat mencakup penilaian hukum dan kemungkinan sertifikasi (sejalan dengan penilaian kesesuaian untuk risiko produk atau standar ISO). Standar ini akan bertindak sebagai insentif bagi pemasok teknologi untuk mencoba memasukkan privasi berdasarkan desain dan kriteria default dalam desain dan penerapan mereka. Beberapa penulis berpendapat bahwa instrumen ini, antara lain, untuk menunjukkan kepatuhan terhadap persyaratan hukum. Yang lain memperingatkan risiko bahwa ketentuan hukum lunak semacam itu dapat meningkatkan ekspektasi dari industri, hingga akhirnya menjadi

persyaratan formal. Dalam hal ini, peran Otoritas Perlindungan Data dan EDPB mungkin krusial dalam memantau kepatuhan dan menilai sertifikasi ini.

Tidak ada instrumen internasional yang mengikat secara hukum yang akan mengatur FRT. Kurangnya regulasi khusus menyebabkan adaptasi terhadap perangkat yang ada, seperti GDPR. Hal ini dapat menyebabkan situasi di mana GDPR ditafsirkan dan/atau diterapkan secara salah. Lebih lanjut, tujuan GDPR bukanlah untuk menjawab tantangan yang ditimbulkan oleh AI, melainkan untuk membangun dasar rezim perlindungan data di Eropa; oleh karena itu, regulasi ini memiliki cakupan yang melampaui AI, dan maknanya dapat hilang atau menyimpang. Saat ini terdapat usulan Undang-Undang AI dari Komisi Eropa, tetapi dengan mempertimbangkan lamanya proses legislasi Uni Eropa, kita masih harus menunggu sebentar untuk melihat instrumen regulasi Uni Eropa di bidang ini.

Untuk mengatasi masalah privasi FRT, mempertimbangkan solusi privasi diferensial dapat bermanfaat. Istilah ini menunjukkan formalisasi matematis dari gagasan sebelumnya bahwa kita harus membandingkan apa yang mungkin dipelajari seseorang dari suatu analisis jika data orang tertentu dimasukkan ke dalam kumpulan data dengan apa yang mungkin dipelajari seseorang jika tidak. Dalam bidang FRT, hal ini memerlukan penerapan 'gangguan' kecil pada templat wajah dan hanya menyimpan templat yang 'terganggu' tersebut, sehingga meminimalkan risiko jika terjadi pelanggaran data.

Alat teknis lainnya, seperti k-anonimisasi, enkripsi homomorfik, atau augmentasi data yang beragam secara fenotip atau demografis menggunakan GAN (Generative Adversarial Networks), mungkin juga memainkan peran penting dalam hal ini. GAN dapat digunakan untuk membuat citra wajah 'buatan' untuk melatih sistem FRT. Kemungkinan ini akan menghasilkan dua keuntungan: di satu sisi, hal ini akan memungkinkan pengenalan wajah yang beragam secara fenotip atau demografis, sehingga mengalahkan bias algoritmik. Di sisi lain, RUU ini akan mempertimbangkan implikasi privasi dan perlindungan data jika terjadi pelanggaran data basis data pelatihan.

Satu pilihan lagi adalah membuat regulasi baru khusus untuk teknologi ini. Namun, opsi ini belum tentu layak, seperti yang ditunjukkan oleh contoh RUU Negara Bagian Washington yang sepenuhnya dikhususkan untuk FRT. RUU ini menandai komitmen regulator terhadap isu keadilan, transparansi, dan privasi FRT, dengan setidaknya mempertahankan akuntabilitas ketika penanganan langsung tidak memungkinkan. 'Laporan akuntabilitas' mungkin menyerupai 'penilaian risiko' yang direnungkan dalam usulan Undang-Undang AI oleh Komisi Uni Eropa. Ini mungkin tampak sebagai formula paling populer ketika berhadapan dengan potensi dampak FRT terhadap hak-hak fundamental. Hanya waktu yang akan membuktikan apakah formula ini berhasil atau tidak.

BAB 3

AI, PEMERINTAH, POLITIK, DAN PSIKOLOGI

3.1 PENDAHULUAN

Perkembangan pesat AI dan implementasinya yang meluas menimbulkan tantangan baru bagi pemerintah dan struktur publik serta organisasi internasional. Salah satu tantangan ini adalah pencegahan kejahatan yang berkaitan dengan penggunaan AI yang jahat (MUAI). Misalnya, undang-undang Rusia tidak menyebutkan kejahatan apa pun yang berkaitan dengan tindakan berbahaya secara sosial melalui penggunaan jaringan saraf, AI, atau tindakan yang dilakukan oleh AI itu sendiri. Tantangan lain muncul dari kurangnya regulasi terkait penggunaan AI untuk memengaruhi keamanan psikologis (PS) suatu masyarakat.

Meskipun contoh Rusia telah disebutkan di atas, hal ini bukan hanya masalah bagi satu negara. Kurangnya regulasi penggunaan AI, termasuk di bidang PS, juga diakui di tingkat Perserikatan Bangsa-Bangsa. Sementara itu, isu MUAI menjadi semakin mendesak ketika membahas masalah destabilisasi psikologis sistem politik, terutama dalam konteks krisis yang disebabkan oleh pandemi COVID-19. Misalnya, selama masa karantina dan isolasi mandiri, jumlah kasus phishing meningkat, dan kerugian yang disebabkan oleh phishing dapat berlipat ganda ketika penipu menggunakan AI. Situs phishing sering kali disamarkan sebagai situs pemerintah, yang dapat menjadi ancaman serius bagi reputasi negara dan badan supranasional. Sampel propaganda teroris yang secara aktif menggunakan gambar AI sudah didistribusikan, dan dalam praktik teroris sudah ada kejahatan yang melibatkan penggunaan kendaraan udara tak berawak (UAV) atau alat pencarian berbasis data besar (untuk melacak korban). Kejahatan semacam itu dapat menjadi faktor yang kuat dalam destabilisasi sistem politik dalam kasus-kasus di mana kekuatan politik yang berlawanan (benar atau salah) menuduh lawan mereka menggunakannya.

Ancaman terpisah dihadirkan oleh kasus-kasus MUAI, di mana penggunaan teknologi AI (seperti deepfake) ditujukan untuk mendestabilisasi PS. Dalam waktu dekat, praktik MUAI mungkin mulai melampaui tidak hanya pengembangan mekanisme hukum untuk mencegah dan menanggapinya, tetapi juga pemahaman realitas baru oleh lembaga-lembaga sipil dan politik, serta oleh komunitas akademis. Oleh karena itu, saat ini perlu mencari cara untuk menanggapi kejahatan-kejahatan di masa depan.

Tugas bab ini adalah untuk mengidentifikasi sejauh mana ancaman terhadap lembaga-lembaga pemerintah dan politik dari MUAI di bidang PS, untuk menentukan status mekanisme hukum untuk mencegah MUAI dan menanggapinya, dan untuk menawarkan rekomendasi-rekomendasi untuk meningkatkan tanggapan ini.

3.2 PENGARUH AI DALAM DINAMIKA POLITIK INDONESIA

Pemanfaatan Artificial Intelligence (AI) dalam politik Indonesia melalui platform Pemilu. AI merupakan contoh nyata bagaimana teknologi cerdas berkontribusi terhadap transformasi strategi kampanye politik berbasis data. Platform ini mengintegrasikan AI dengan

Big Data analytics untuk mengolah beragam informasi pemilih, mencakup aspek demografis, kecenderungan preferensi politik, pola partisipasi elektoral, serta dinamika opini publik yang berkembang di media sosial. Kemampuan AI dalam mengidentifikasi pola dan tren tersembunyi dari data skala besar memungkinkan tim kampanye menyusun strategi komunikasi yang jauh lebih personal dan terarah. Misalnya, pesan politik dapat disesuaikan dengan kebutuhan spesifik kelompok masyarakat tertentu, baik berdasarkan wilayah, usia, maupun latar belakang sosial-ekonomi. Di samping itu, AI pada Pemilu. AI juga berfungsi sebagai sistem peringatan dini (*early warning system*) terhadap isu-isu politik yang berpotensi memengaruhi citra kandidat. Dengan mendeteksi perubahan sentimen publik secara real-time, kandidat maupun partai politik memiliki peluang lebih besar untuk segera melakukan penyesuaian narasi, mengoreksi kesalahan strategi, atau memperkuat isu yang sedang positif di mata publik. Hal ini tidak hanya meningkatkan efektivitas komunikasi politik, tetapi juga efisiensi penggunaan sumber daya kampanye.

Lebih jauh, pengaruh AI dalam politik Indonesia melalui pemilu dapat dilihat dari kontribusinya dalam mendorong kampanye berbasis bukti (*evidence-based campaign*), di mana setiap keputusan strategis didasarkan pada analisis data yang objektif, bukan semata intuisi politik. Dengan demikian, peluang keberhasilan kampanye dapat ditingkatkan melalui optimasi alokasi sumber daya, penyusunan agenda politik yang lebih relevan dengan kebutuhan masyarakat, serta peningkatan keterhubungan antara kandidat dan pemilih. Meski demikian, literatur juga menegaskan adanya implikasi etis dan risiko yang perlu diwaspadai. Penggunaan AI dalam politik rawan menimbulkan masalah privasi terkait pengumpulan dan pemrosesan data pribadi pemilih, serta membuka potensi manipulasi opini publik melalui algoritme yang disalahgunakan. Oleh karena itu, keberhasilan penerapan AI dalam politik Indonesia harus diimbangi dengan regulasi yang ketat, transparansi penggunaan data, dan mekanisme akuntabilitas yang jelas, agar inovasi teknologi ini dapat berfungsi sebagai instrumen demokratisasi, bukan sebagai alat manipulasi politik.

3.3 ANCAMAN PSIKOLOGIS BARU TERHADAP PEMERINTAH DAN INSTITUSI POLITIK

Ancaman terhadap PS dari MUAI dapat dibagi menjadi tiga tingkat. Pada tingkat pertama, ancaman ini terkait dengan interpretasi yang sengaja diputarbalikkan tentang keadaan dan konsekuensi pengembangan AI untuk kepentingan kelompok antisosial. Dalam hal ini, AI sendiri tidak terlibat dalam destabilisasi keamanan psikologis internasional (IPS). Dampak destruktifnya (terbuka atau tersembunyi) adalah terciptanya citra palsu AI di benak masyarakat.

Bidang MUAI sangat luas: penggunaan drone yang tidak dapat dibenarkan; ancaman serangan siber terhadap infrastruktur yang rentan; manipulasi mata uang kripto; dan masih banyak lagi. Jika MUAI tidak ditujukan untuk mengelola target audiens di ranah psikologis tetapi terutama untuk melakukan tindakan jahat lainnya (misalnya, penghancuran infrastruktur penting), kita dapat membahas tingkat kedua MUAI.

Penggunaan sarana dan metode perang psikologis secara profesional dapat meningkatkan persepsi tingkat ancaman di atas atau di bawah tingkat yang wajar. Terlebih lagi,

penggunaan AI dalam perang psikologis justru membuat kampanye manajemen persepsi yang tersembunyi (laten) menjadi lebih berbahaya; hal ini akan semakin meningkat di masa mendatang. Oleh karena itu, MUAI, yang utamanya ditujukan untuk menimbulkan kerusakan di ranah psikologis, patut mendapat perhatian yang independen dan sangat cermat, karena merupakan ancaman tingkat ketiga yang khusus bagi IPS.

MUAI dapat melibatkan beragam teknologi AI. Lebih dari sebelumnya, perkembangan AI yang pesat di dunia merupakan masalah nyata bagi kerentanan objek fisik yang dikendalikan oleh sistem kompleks yang sebagian atau seluruhnya berbasis AI. Integrasi aktif internet industri, tempat produk AI beroperasi, dan penggunaan analitik prediktif dalam industri, transportasi, dan pertanian, dengan tetap memastikan tingkat keamanan yang memadai, dapat meningkatkan efisiensi perusahaan, industri, dan kemudian perekonomian secara signifikan. Namun, memasukkan data yang terdistorsi atau sepenuhnya palsu ke dalam AI yang belajar mandiri, misalnya, dapat menyebabkan penyemprotan bahan kimia beracun di ladang atau menciptakan kemacetan lalu lintas yang parah atau kecelakaan (menggunakan mekanisme prediksi intensitas lalu lintas). Jenis MUAI ini dapat mendiskreditkan mekanisme AI itu sendiri dan produsennya, serta mereka yang membuat keputusan untuk mengimplementasikannya, dan bahkan pengguna biasa yang memilih produk AI tertentu.

Risiko baru tidak diragukan lagi terkait dengan implementasi produk AI sintetis, yang mencakup, misalnya, pengelolaan objek fisik beserta sistem identifikasi pengguna berdasarkan gambar suara atau video. Masalah ini berkaitan dengan 'penggunaan deepfake yang jahat'. Dalam arti sempit, proses pembuatan deepfake berarti penambahan satu gambar atau video digital di atas gambar atau video lain sedemikian rupa sehingga gambar yang ditambahkan tampak seperti bagian dari gambar asli. Namun, tanpa mengabaikan perbedaan dalam teknologi tertentu, kami menganggap istilah 'deepfake' dapat digunakan dalam arti yang lebih luas, menggabungkan proses asli dengan serangkaian teknologi yang ada dan yang akan datang untuk membangun pseudo-realitas. Teknologi-teknologi ini didasarkan pada kemampuan AI untuk membuat atau memodifikasi gambar, video, suara, dan teks.

Penggunaan chatbot yang jahat dengan tambahan sistem sintesis yang mampu menggunakan suara orang sungguhan jauh lebih berbahaya daripada penggunaan chatbot biasa, yang kurang meyakinkan dalam percakapan. Jika saat ini chatbot dengan suara Robbie Williams dibuat untuk tujuan yang baik,⁶ tidak ada yang akan mencegah penjahat di masa depan untuk membuat bot semacam itu dengan suara Osama bin Laden atau orang semacam itu. Di bidang informasi, terdapat juga teknologi 'peringkat dan penurunan peringkat'. Teknologi pemeringkatan menggunakan AI untuk meningkatkan, sementara teknologi penurunan peringkat digunakan untuk menurunkan, peringkat situs di mesin pencari, dengan menyesuaikan popularitas suatu peristiwa atau fenomena. Teknologi penurunan peringkat memungkinkan, misalnya, perusahaan internet, jika mereka terlibat dalam perang psikologis di pihak kekuatan politik tertentu, untuk menggunakan mesin pencari mereka untuk mendiskreditkan lawan politik, dan sudah ada bukti nyata mengenai hal ini.

Diskreditasi politik terhadap struktur pemerintah 'berkontribusi' pada penyebaran phishing, di mana penjahat mendapatkan akses ke data pengguna rahasia (nama pengguna

dan kata sandi) dengan mengirimkan email massal atas nama merek populer, atau pesan pribadi dari berbagai layanan, misalnya, atas nama bank atau jejaring sosial. Phishing AI bisa lebih sulit diidentifikasi daripada serangan phishing biasa, karena personalisasi pesan.

Singkatnya, adalah mungkin untuk menentukan tujuan MUAI dilakukan dalam rangka perang psikologis. Perang psikologis selalu ditujukan untuk memberikan pukulan langsung (meskipun seringkali laten) terhadap kesadaran publik dan (melalui kemenangan di bidang ini) kemenangan umum tanpa pertumpahan darah atas musuh. MUAI dalam konteks IPS ditujukan untuk mendapatkan keuntungan dalam perang psikologis melalui bentuk-bentuk pengaruh baru secara kuantitatif dan kualitatif pada kesadaran individu dan publik.

3.4 KLASIFIKASI MUAI

Adalah mungkin untuk menawarkan klasifikasi MUAI berikut menurut tingkat kesiapannya: (1) praktik MUAI yang ada; (2) kapabilitas MUAI yang ada namun belum digunakan dalam praktik – kapabilitas ini terkait dengan beragam fitur AI baru yang berkembang pesat, yang tidak semuanya langsung tercakup dalam rangkaian fitur MUAI yang telah diimplementasikan; (3) peluang MUAI di masa mendatang berdasarkan perkembangan yang sedang berlangsung dan penelitian di masa mendatang (penilaian harus diberikan untuk jangka pendek, menengah, dan panjang); dan (4) risiko yang belum teridentifikasi, juga dikenal sebagai 'unknown unknowns'. Tidak semua perkembangan di bidang AI dapat dinilai secara akurat. Kesiapan untuk menghadapi risiko laten MUAI yang tidak terduga sangatlah penting.

Klasifikasi MUAI lainnya juga dimungkinkan: (1) berdasarkan cakupan teritorial (lokal, regional, global); (2) berdasarkan tingkat kerusakan (kecil, signifikan, besar, katastrofik); (3) berdasarkan kecepatan propagasi (lambat, cepat, cepat); dan (4) berdasarkan bentuk propagasi (terbuka, laten).

Risiko MUAI meningkat secara signifikan akibat penggunaan teknologi AI yang terintegrasi oleh para penyusup. Namun, kita tidak boleh lupa bahwa ancaman AI terhadap PS terutama bersifat antropogenik – disebabkan oleh aktivitas manusia, bukan hal-hal spesifik AI (setidaknya pada tahap pengembangan AI saat ini).

3.5 TANGGAPAN HUKUM ATAS MUAI DI UNI EROPA DAN AS

Pengembangan sarana hukum untuk penuntutan MUAI menjadi salah satu tugas utama para legislator yang sedang mengembangkan mekanisme untuk mengatur AI. Berdasarkan dokumen hukum dan kebijakan Uni Eropa dan Amerika Serikat (yang saat ini memimpin dalam bidang regulasi AI), perubahan sikap pemerintah terhadap masalah MUAI di bidang PS dapat ditelusuri.

UNI EROPA

Uni Eropa diakui sebagai salah satu yurisdiksi utama di mana regulasi robotika dan AI sedang dikembangkan. Resolusi Parlemen Eropa tanggal 16 Februari 2017 dengan Rekomendasi kepada Komisi Aturan Hukum Perdana tentang Robotika, yang pertama kali mengusulkan istilah 'orang elektronik', merupakan hal yang inovatif. Aspek etika dari resolusi ini terutama merupakan indikasi dari perpecahan yang semakin besar dalam masyarakat dan

menyusutnya kelas menengah, yang mana perkembangan robotika dapat menyebabkan konsentrasi kekayaan dan pengaruh yang tinggi di tangan minoritas. Tak kalah pentingnya, karena alasan ini, Parlemen Eropa menyerukan kepada mereka yang terlibat 'dalam pengembangan dan komersialisasi aplikasi AI' untuk 'membangun keamanan dan etika sejak awal'. Dalam aspek memastikan PS, sangat penting untuk menekankan bahaya konsentrasi AI di tangan elit negara anti-nasional, serta aktor non-negara agresif lainnya, termasuk organisasi kriminal (hingga teroris).

Pendekatan hak asasi manusia terhadap masalah MUAI dinyatakan dalam publikasi terkait dari Badan Kerja Sama Penegakan Hukum Uni Eropa (Europol). Europol sedang mempertimbangkan berbagai ancaman yang terkait dengan MUAI, menyadari bahwa teknologi AI dan pembelajaran mesin terus mengubah lanskap keamanan dan menjadi lebih mudah diakses. Berdasarkan analisis masalah MUAI yang dilakukan di Universitas Oxford, Europol merekomendasikan agar lembaga penegak hukum mempelajari berbagai skenario yang mungkin terjadi. Misalnya, penyerang dapat menggunakan AI untuk membuat dan mendistribusikan konten berbahaya (serangan rekayasa sosial atau email phishing), dan untuk mendeteksi vektor serangan baru atau kerentanan calon korban. Teknologi AI dapat digunakan untuk memilih target serangan, menetapkan prioritas, atau menanggapi perubahan perilaku target. Dampak disinformasi skala besar terhadap warga negara dicatat sebagai masalah serius bagi UE. Penggunaan AI dapat semakin meningkatkan dampak ancaman hibrida, karena produksi massal informasi palsu dapat diotomatisasi oleh penyerang dengan sedikit atau tanpa pengetahuan teknis. Teknologi deepfake memberi penyerang kemampuan untuk menyebarkan misinformasi dengan menyamar sebagai orang lain. Penjahat telah menggunakan deepfake suara, menyamar sebagai eksekutif perusahaan dalam upaya untuk menipu karyawan organisasi. Semua ancaman ini dapat menjadi subjek hukum dalam undang-undang UE yang diperluas, terutama karena proses peningkatan ini didukung oleh dokumen strategis baru di bidang AI.

Komisi Eropa merilis rencananya untuk masa depan AI di UE pada 19 Februari 2020. Strategi tersebut ditetapkan dalam dua dokumen utama. Yaitu 'Laporan tentang Implikasi Keselamatan dan Tanggung Jawab Kecerdasan Buatan, Internet of Things, dan Robotika' dan 'Buku Putih tentang Kecerdasan Buatan - Pendekatan Eropa terhadap Keunggulan dan Kepercayaan'. Yang terakhir berisi bagian berjudul 'Ekosistem Kepercayaan: Kerangka Kerja Regulasi untuk AI', yang mengevaluasi alat yang tersedia bagi Uni untuk mengatur teknologi. Komisi UE mencatat, 'Warga negara takut tidak berdaya dalam membela hak dan keselamatan mereka ketika menghadapi asimetri informasi dalam pengambilan keputusan algoritmik, dan perusahaan khawatir dengan ketidakpastian hukum'. Sangat penting bagi Komisi untuk mencatat kemampuan AI untuk melindungi keamanan warga negara dan memungkinkan mereka menikmati hak-hak dasar mereka, dan risikonya, termasuk yang disebabkan oleh MUAI. 'Kurangnya kepercayaan' dicatat sebagai faktor utama yang menghambat penyebaran AI yang lebih luas. Dengan demikian, strategi Uni Eropa membawa kita pada asumsi bahwa, dengan latar belakang perang informasi yang terjadi di seluruh dunia, ancaman tingkat pertama terhadap PS (keinginan aktor agresif untuk memanfaatkan ketakutan dan

ketidakpercayaan warga negara guna mendiskreditkan politisi yang secara aktif mengadvokasi penerapan AI) dapat menjadi sangat relevan.

Menurut Komisi Eropa, penggunaan AI dapat menyebabkan pelanggaran hak berkumpul secara bebas, diskriminasi berdasarkan gender, ras atau etnis, agama atau kepercayaan, disabilitas, usia atau orientasi seksual, serta pelanggaran hak atas perlindungan yang efektif dan pengadilan yang adil, dan hak-hak konsumen. Pada saat yang sama, Komisi terutama menyoroti sifat antropogenik dari semua pelanggaran yang disebutkan di atas, yang juga dapat menjadi titik awal untuk mengembangkan standar pencegahan dan penuntutan MUI.

Komisi Eropa pada April 2021 merilis 'Proposal Regulasi tentang Pendekatan Eropa untuk Kecerdasan Buatan'. Proposal tersebut mencatat bahwa

Teknologi AI juga dapat disalahgunakan dan menyediakan alat yang baru dan ampuh untuk praktik manipulatif, eksploitatif, dan kontrol sosial. Praktik-praktik tersebut sangat berbahaya dan harus dilarang karena bertentangan dengan nilai-nilai Uni Eropa tentang penghormatan terhadap martabat manusia, kebebasan, kesetaraan, demokrasi, dan supremasi hukum serta hak-hak dasar Uni Eropa, termasuk hak atas non-diskriminasi, perlindungan data dan privasi, serta hak-hak anak.

Regulasi hukum AI secara umum, dan MUI khususnya, masih dalam tahap awal perkembangannya, bahkan di Uni Eropa. Hal ini ditunjukkan oleh contoh-contoh regulasi deepfake. Di berbagai negara, baru langkah legislatif pertama yang diambil di bidang ini. Pada tahun 2018, Uni Eropa merilis dokumen 'Menangani Disinformasi Daring: Pendekatan Eropa', yang mengkaji disinformasi secara umum dan menyediakan serangkaian pedoman, termasuk pedoman untuk perlindungan terhadap deepfake.

Disinformasi adalah alat pengaruh yang ampuh dan murah – dan seringkali menguntungkan secara ekonomi – ... teknologi baru, terjangkau, dan mudah digunakan kini tersedia untuk membuat gambar palsu dan konten audio-visual (disebut 'deep fake'), menawarkan cara yang lebih ampuh untuk memanipulasi opini publik.

Uni Eropa menyerukan partisipasi publik yang lebih besar dan transparansi yang lebih besar dalam mengidentifikasi asal konten internet, dan jaringan Eropa independen telah dibentuk untuk memverifikasi sumber dan proses pembuatannya.

AMERIKA SERIKAT

Di AS, isu keamanan AI dibahas dalam dokumen berjudul 'Rencana Strategis Penelitian dan Pengembangan Kecerdasan Buatan Nasional 2016'. Hal ini menunjukkan bahwa masalah MUI dapat dikaitkan dengan dua strategi. Yang pertama adalah 'memastikan keselamatan dan keamanan sistem AI', yang menyiratkan memastikan tidak hanya keamanan proses produksi produk AI, tetapi juga kejelasan pekerjaannya di setiap tahap. Strategi ini juga bertujuan untuk memberikan perlindungan yang andal bagi produk AI dari serangan siber.

Strategi untuk mengembangkan AI yang 'taat hukum' ('memahami dan menangani implikasi etika, hukum, dan sosial dari AI') juga berkaitan erat dengan tugas mempertahankan kendali manusia atas AI. Para penyusun Rencana Strategis menekankan pentingnya kemampuan AI untuk membuat keputusan independen sesuai dengan standar etika manusia dan peraturan perundang-undangan yang berlaku. Dalam hal ini, sangat penting untuk menunjukkan perlunya tim peneliti interdisipliner, setidaknya untuk bekerja menjawab pertanyaan tentang data mana yang akan digunakan untuk melatih AI. Dalam kasus MUIAI, termasuk MUIAI yang bertujuan untuk mendestabilisasi PS, pendekatan ini merupakan titik awal yang tepat untuk membahas subjektivitas pelanggaran di masa mendatang.

Inisiatif legislatif di bidang AI sedang berkembang di Amerika Serikat. Pada tahun 2018, lima rancangan undang-undang diusulkan untuk mengatur AI, pada tahun 2019 ada delapan, dan selama paruh pertama tahun 2020 (hingga 19 Juni) ada sepuluh. Teknologi deepfake adalah salah satu teknologi MUIAI yang bertujuan untuk mendestabilisasi PS yang paling aktif dibahas oleh para legislator Amerika. Proyek pertama adalah rancangan undang-undang federal untuk mengatur deepfake, 'Malicious Deep Fake Prohibition Act of 2018', yang diperkenalkan di Amerika Serikat pada bulan Desember 2018, dan 'Deep Fakes Accountability Act' diperkenalkan pada bulan Juni 2019. Undang-undang yang menentang deepfake juga telah diperkenalkan di beberapa negara bagian, termasuk California, New York, dan Texas.

Sebagai hasil dari diskusi yang intens tentang isu-isu MUIAI di Amerika Serikat pada tanggal 20 Desember 2019, Presiden Donald Trump menandatangani undang-undang federal pertama di negara itu yang terkait dengan 'deepfake'. Undang-undang deepfake merupakan bagian dari Undang-Undang Otorisasi Pertahanan Nasional untuk Tahun Anggaran 2020 (NDAA). Selain deepfake, bagian terpisah dari undang-undang tersebut (bagian 5708) didedikasikan untuk teknologi pengenalan wajah dan penggunaannya oleh komunitas intelijen. Ini juga merupakan langkah penting dalam mencegah serangan MUIAI terhadap keamanan nasional negara, karena undang-undang tersebut melarang transfer teknologi kepada aktor politik yang agresif.

Berdasarkan undang-undang tersebut, Direktur Intelijen Nasional (DNI) harus, selambat-lambatnya satu tahun setelah diadopsi, memberikan laporan kepada Kongres tentang bagaimana teknologi pengenalan wajah digunakan untuk tujuan intelijen, dalam keadaan apa teknologi tersebut harus dibagikan dengan entitas di luar yurisdiksi Amerika Serikat, dan apakah penggunaan teknologi ini mengancam hak konstitusional warga negara.

Namun, perselisihan mengenai teknologi tersebut belum mereda dengan diadopsinya undang-undang tersebut. Pada tahun 2020, para legislator di Dewan Perwakilan Rakyat yang "waspada terhadap pengenalan wajah" sedang menyusun undang-undang yang akan menghentikan perkembangannya hingga potensi risikonya dapat dipahami dan dimitigasi. Pada bulan Juni 2020, sekelompok legislator di Amerika Serikat kembali mengusulkan undang-undang yang akan memberlakukan moratorium federal atas penggunaan teknologi pengenalan wajah oleh lembaga penegak hukum, upaya pertama untuk memberlakukan larangan sementara terhadap teknologi tersebut di seluruh negeri. Moratorium federal yang diusulkan ini muncul setelah kota-kota seperti San Francisco dan Boston, dengan alasan

masalah privasi dan bias rasial dalam teknologi tersebut, mengesahkan larangan mereka sendiri terhadap pengenalan wajah.

Tidak seperti laporan tentang teknologi pengenalan wajah, laporan tentang dampak teknologi deepfake pada keamanan nasional AS harus diserahkan ke Kongres paling lambat 180 hari setelah undang-undang disahkan. NDAA mendefinisikan deepfake sebagai 'media yang dimanipulasi mesin', yang merupakan definisi yang lebih luas daripada yang ada dalam RUU 2018. Interpretasi deepfake yang lebih luas ini tidak hanya mencakup video, tetapi juga materi lain yang dibuat dengan bantuan AI dan dirancang untuk membentuk pseudo-realitas. Efektivitas transisi ke interpretasi ini telah terbukti dalam praktik nyata, seperti dalam percobaan yang dilakukan oleh Max Weiss dari Harvard. Dia menggunakan program pembangkit teks untuk membuat 1.000 komentar sebagai tanggapan atas panggilan pemerintah tentang masalah Medicaid. 'Semua komentar ini unik, dan terdengar seperti orang sungguhan yang mengadvokasi posisi kebijakan tertentu. Mereka menipu administrator Medicaid.gov, yang menerimanya sebagai kekhawatiran asli dari manusia yang sebenarnya'. Kemudian, Weiss mengidentifikasi komentar-komentar tersebut dan meminta agar komentar-komentar tersebut dihapus, 'agar tidak ada perdebatan kebijakan yang bias secara tidak adil'. Eksperimen ini sangat instruktif dan menunjukkan kemungkinan praktik MUAI melalui pembuatan teks palsu, yang belum kita ketahui.

Aktivitas legislatif dan adopsi dokumen perencanaan strategis di Uni Eropa dan Amerika Serikat memungkinkan kita untuk melacak perkembangan regulasi AI seiring dengan identifikasi isu-isu AI secara bertahap, termasuk di bidang PS. Preseden di bidang ini seringkali mendorong peningkatan mekanisme hukum untuk penuntutan MUAI, tetapi saat ini peran penting pendekatan interdisipliner dalam memprediksi kejahatan di masa depan sudah terlihat jelas (lebih dari sebelumnya, karena pesatnya perkembangan teknologi).

3.6 TANGGAPAN HUKUM DI TIONGKOK DAN RUSIA TIONGKOK

Menurut sumber terbuka yang tersedia, regulasi AI di Tiongkok tampaknya masih dalam tahap awal pengembangan. Dokumen resmi seperti 'Pemberitahuan Dewan Negara yang Mengeluarkan Rencana Pengembangan Kecerdasan Buatan Generasi Baru' berfokus terutama pada aspek positif penggunaan AI – AI membantu dalam membuat penilaian dan prediksi yang akurat tentang tren utama dalam pengembangan infrastruktur dan jaminan sosial. Di bidang jaminan sosial, pemerintah Tiongkok juga melihat peluang positif yang luas untuk 'memahami perubahan kesadaran dan psikologi kelompok secara tepat waktu, merespons pengambilan keputusan secara aktif, meningkatkan kemampuan dan tingkat tata kelola sosial secara signifikan, dan sangat diperlukan untuk pemeliharaan stabilitas sosial yang efektif'. Pada saat yang sama, tantangan baru diakui, termasuk bahaya dampak destruktif AI terhadap 'manajemen pemerintahan, keamanan ekonomi dan stabilitas sosial, dan bahkan tata kelola global', yang dapat menyebabkan 'masalah perubahan struktur ketenagakerjaan, dampak hukum dan etika sosial, pelanggaran privasi pribadi, dan tantangan hubungan internasional'.

Pada tahun 2019, sebuah komite ahli yang dibentuk oleh Kementerian Sains dan Teknologi Tiongkok (MOST) merilis sebuah dokumen yang menguraikan delapan prinsip tata kelola AI dan 'AI yang bertanggung jawab'.

Prinsip pertama, 'Harmoni dan Keramahan', menyiratkan bahwa 'pengembangan AI ... harus sesuai dengan nilai-nilai kemanusiaan, etika, dan moralitas ... harus didasarkan pada premis menjaga keamanan masyarakat dan menghormati hak asasi manusia, menghindari penyalahgunaan, dan melarang penyalahgunaan dan penerapan yang jahat'. Dengan demikian, di Tiongkok, isu penanggulangan MUAI, termasuk di bidang PS, secara resmi diakui sebagai prioritas.

Isu-isu MUAI banyak dibahas di kalangan akademis dan publik, yang memberikan banyak peluang untuk mengembangkan mekanisme regulasi. Sekelompok peneliti dari Universitas Tsinghua telah menyoroti berbagai ancaman yang terkait dengan MUAI. Secara khusus, mereka menunjukkan bahwa ruang/penyimpanan daring virtual memfasilitasi pengumpulan dan pertukaran data pribadi, analisis dan pertukaran informasi (termasuk data identifikasi, informasi medis, catatan kredit, lokasi pribadi, dan informasi pergerakan). Namun, pada saat yang sama, hal ini menyulitkan untuk menentukan penyebab dan tingkat kebocoran data tersebut. Dari sudut pandang PS, perlu diperhatikan kondisi-kondisi yang akan memengaruhi jiwa manusia, seperti 'ketidakpastian dan ketidakterbalikan dari pengaburan waktu dan ruang', serta realitas virtual dan objektif, yang menurut para ahli Tiongkok berpotensi menimbulkan risiko.

Di antara ancaman utama MUAI yang dicatat oleh laporan Universitas Tsinghua adalah ancaman yang akan kami klasifikasikan sebagai milik tingkat kedua: penipuan (termasuk penipuan di jejaring sosial, berdasarkan informasi pribadi yang diperoleh secara ilegal), peretasan sistem otentikasi berbasis AI, dan penggunaan jahat UAV, kendaraan otonom dan robot yang dilengkapi dengan AI. Masalah keamanan psikologis dibahas dalam struktur tertutup, terutama militer, tetapi perlu ditekankan bahwa penting bahwa masalah ini dipelajari oleh pusat analisis sipil yang menangani masalah implementasi AI, terutama dalam konteks pengembangan inisiatif keamanan siber Tiongkok.

Contoh paling terkenal dari regulasi MUAI di Tiongkok adalah larangan penggunaan deepfake untuk menyesatkan khalayak. Pada tahun 2019, Tiongkok mengumumkan aturan baru yang mengatur konten video dan audio di internet, termasuk larangan publikasi dan distribusi 'berita palsu' yang dibuat menggunakan AI dan realitas virtual. Teknologi deepfake dapat "membahayakan keamanan nasional, mengganggu stabilitas sosial, mengganggu ketertiban sosial, dan melanggar hak serta kepentingan sah orang lain", menurut transkrip konferensi pers yang dipublikasikan di situs web Administrasi Ruang Siber Tiongkok (CAC). Setiap penggunaan teknologi ini harus ditandai dengan jelas dan mudah dilihat oleh pengguna internet. Kegagalan untuk mematuhi aturan ini dapat dianggap sebagai tindak pidana, menurut situs web CAC. Perlu dicatat bahwa bukan teknologinya yang dilarang, melainkan penyesatan yang disengaja terhadap audiens dengan bantuan teknologi tersebut.

Rusia

MUAI menjadi sorotan di Rusia, sebagaimana dibuktikan oleh pernyataan Presiden Vladimir Putin pada September 2017. "Kecerdasan buatan bukan hanya masa depan Rusia, melainkan masa depan seluruh umat manusia. Ada peluang dan ancaman besar yang sulit diprediksi saat ini." Pada tahun 2018, Kementerian Pertahanan Rusia menyelenggarakan konferensi pertama bertajuk "Kecerdasan Buatan: Masalah dan Solusi". Dalam konferensi ini, Menteri Pertahanan Sergei Shoigu menugaskan para spesialis militer dan sipil untuk bergabung dalam mengembangkan teknologi AI guna melawan potensi ancaman di bidang keamanan teknologi dan ekonomi. Dengan demikian, fokus utama perhatian tertuju pada ancaman MUAI yang tidak secara langsung berkaitan dengan dampak psikologis, tetapi dapat "beregenerasi" menjadi ancaman keamanan tingkat kedua.

Pada 10 Oktober 2019, Rusia mengadopsi 'Strategi Nasional Pengembangan Kecerdasan Buatan hingga 2030'. Poin-poin mengenai hak asasi manusia, yang dapat menjadi titik awal pengembangan mekanisme pencegahan dan penuntutan MUAI, terdapat dalam dua bagian dokumen, salah satunya adalah 'Prinsip-prinsip dasar pengembangan dan penggunaan teknologi kecerdasan buatan' (bagian III). Di antara prinsip-prinsip tersebut, yang harus diperhatikan dalam penerapan Strategi, adalah:

- a) perlindungan hak asasi manusia dan kebebasan: memastikan perlindungan hak asasi manusia dan kebebasan yang dijamin oleh undang-undang Rusia dan internasional, termasuk hak untuk bekerja, dan memberikan warga negara kesempatan untuk memperoleh pengetahuan dan keterampilan agar berhasil beradaptasi dengan ekonomi digital;
- b) keamanan: tidak dapat diterimanya penggunaan AI untuk tujuan yang secara sengaja merugikan warga negara dan badan hukum, serta mencegah dan meminimalkan risiko akibat negatif dari penggunaan teknologi AI.

Salah satu tujuan utama penciptaan sistem komprehensif untuk mengatur hubungan masyarakat yang timbul dari pengembangan dan implementasi teknologi AI dalam strategi Rusia adalah pengembangan aturan etika untuk interaksi manusia dengan AI.

Undang-undang federal berjudul 'Tentang Pelaksanaan Eksperimen untuk Menetapkan Regulasi Khusus dalam Rangka Menciptakan Kondisi yang Diperlukan untuk Pengembangan dan Implementasi Teknologi AI di Subjek Federasi Rusia – Kota Penting Federal Moskow' memicu reaksi positif dan negatif di media. Undang-undang tersebut memuat indikasi bahwa eksperimen tidak boleh melanggar hak-hak warga negara:

Hasil dari pembentukan rezim hukum percontohan tidak boleh berupa pembatasan hak dan kebebasan konstitusional warga negara, pengenaan tugas tambahan, pelanggaran kesatuan ruang ekonomi di wilayah Federasi Rusia, atau penyusutan lain terhadap perlindungan yang menjamin perlindungan hak-hak warga negara dan badan hukum.

Tugas eksperimental untuk membangun rezim hukum sebelumnya telah ditetapkan dalam Strategi Nasional untuk Pengembangan AI. Saat ini, mengingat sikap ambigu terhadap eksperimen ini di masyarakat Rusia, lembaga pemerintah kemungkinan besar tidak hanya

harus mengendalikan eksperimen tersebut, tetapi juga menjelaskan sifat teknologi yang diperkenalkan, serta manfaat dan risiko penggunaannya bagi warga negara.

Di Rusia, seperti di banyak negara di dunia, tidak ada peraturan khusus untuk penggunaan set data besar oleh robot yang dilengkapi dengan AI, yang dapat memengaruhi penerapan rezim hukum eksperimental. Mengingat undang-undang Rusia di bidang AI saat ini masih dalam tahap awal pengembangan, masih banyak pekerjaan yang harus dilakukan. Bagaimana undang-undang tersebut akan menghormati prinsip transparansi, membuat karya AI dapat dipahami, dan menetapkan dalam kebijakan dan hukum federal akses non-diskriminatif bagi warga negara terhadap hasil AI? Ini adalah poin-poin mendasar dari sudut pandang hak-hak warga negara, dan sikap warga negara itu sendiri terhadap implementasi inovasi. Saat ini, poin-poin tersebut sangat penting bagi perkembangan masyarakat dan teknologi.

3.7 PENDEKATAN SISTEMIK BAGI KEBERHASILAN HUKUM BALASAN

Tidak diragukan lagi, pendekatan interdisipliner sistemik dan kolaborasi antara peneliti dan praktisi sangat penting dalam pengembangan norma hukum dan rencana kebijakan strategis di bidang AI. Pengalaman Eropa dalam mempromosikan penelitian dan berbagi keahlian, yang dijelaskan secara rinci di sini, mungkin bermanfaat dalam hal ini. Untuk memantau implementasi AI, Komisi Eropa telah membentuk AI Watch, sebuah layanan informasi untuk memantau perkembangan, implementasi, dan dampak AI di Eropa. AI Watch merupakan bagian dari Pusat Penelitian Gabungan (JRC) Komisi Eropa, bekerja sama dengan Direktorat Jenderal Jaringan Komunikasi, Konten, dan Teknologi (DG CONNECT). Komisi Eropa, khususnya, bekerja sama dengan OECD dalam analisis strategi AI nasional Negara-negara Anggota UE. OECD dan Komisi Eropa berbagi konten untuk dipublikasikan di platform mereka, seperti OECD AI Policy Observatory dan AI Watch. Tugas penting AI Watch adalah pembentukan komunitas riset di bidang AI di Uni Eropa.

Pengalaman kerja sama internasional yang dibangun oleh negara-negara BRICS (Brasil, Rusia, India, Tiongkok, dan Afrika Selatan) merupakan contoh yang bermanfaat. BRICS menciptakan platform untuk pertukaran ide dan hasil riset, CyberBRICS. Para pengembang proyek CyberBRICS menetapkan tiga tujuan: membandingkan regulasi yang ada; mengidentifikasi praktik terbaik; dan mengembangkan proposal kebijakan di bidang regulasi keamanan siber (termasuk regulasi data pribadi), kebijakan akses internet, dan strategi digitalisasi otoritas publik di negara-negara BRICS. Pencapaian tujuan-tujuan ini akan membantu mengembangkan mekanisme hukum dan kebijakan untuk mengatur TIK, khususnya AI.

Para menteri BRICS yang bertanggung jawab atas sains, teknologi, dan inovasi telah bekerja sama sejak 2014, mengadopsi sejumlah dokumen untuk mengembangkan kerangka hukum, termasuk 'Nota Kesepahaman tentang Kerja Sama dalam Sains, Teknologi, dan Inovasi'. Sejak itu, kerja sama di bidang TIK telah berkembang pesat, sebagaimana dibuktikan oleh promosi Kemitraan Digital BRICS pada tahun 2016, adopsi Deklarasi KTT Presiden BRICS untuk Kolaborasi untuk Pertumbuhan Inklusif dan Kemakmuran Bersama dalam Revolusi

Industri ke-4, dan pengembangan Kerangka Kerja Pendukung untuk Jaringan Inovasi BRICS. Kita harus mempertimbangkan inisiatif baru seperti Kemitraan BRICS untuk Revolusi Industri Baru (PartNIR), Jaringan Inovasi BRICS (Jaringan iBRICS), dan Institut Jaringan Masa Depan BRICS. Pendirian Pusat Transfer Teknologi BRICS pertama (di Kunming) dan Institut Jaringan Masa Depan BRICS pertama (di Shenzhen) di Tiongkok tidak hanya menunjukkan kepemimpinan negara tersebut dalam AI, tetapi juga minatnya dalam mengembangkan kerja sama teknologi di BRICS. Liu Duo, Presiden Akademi Teknologi Informasi dan Komunikasi Tiongkok, telah mengatakan bahwa Institut di Shenzhen akan berfokus pada penelitian kebijakan tentang 5G, internet industri, AI, internet kendaraan, dan teknologi lainnya. Upaya tambahan akan dilakukan untuk mendorong pertukaran ide antara kelima negara dan untuk menyelenggarakan lebih banyak acara pelatihan dan program pertukaran. Platform iBRICS didirikan oleh KTT 2019 di Brasil dan menyediakan kontak antara taman sains dan kluster teknologi negara-negara BRICS, asosiasi khusus, dan struktur pendukung untuk mendorong perusahaan rintisan, termasuk di bidang AI.

Sayangnya, cakupan bab ini yang terbatas tidak memungkinkan kami untuk menggambarkan lebih banyak pengalaman praktis di bidang regulasi MUI yang masih baru. Namun, sudah jelas bahwa akses terbuka terhadap dokumen perencanaan strategis memungkinkan warga negara (termasuk peneliti) untuk mendapatkan gambaran lengkap tentang perkembangan AI, serta risiko dan ancaman yang terkait dengannya. Saat ini, terdapat diskusi aktif mengenai prinsip-prinsip dasar yang seharusnya mengatur penggunaan AI dalam pengumpulan, pemrosesan, dan penggunaan data. Dokumen-dokumen strategis Uni Eropa menyatakan perlunya memastikan kontrol pribadi oleh warga negara atas data pribadi mereka. Di Amerika Serikat, misalnya, sebelumnya terdapat usulan untuk merevisi prinsip ini demi pengelolaan data pribadi yang efektif demi kepentingan warga negara. Namun, dalam kedua kasus tersebut, penggunaan AI tidak dapat dilindungi dari ancaman yang bersifat antropogenik, ketika aktor politik atau bisnis yang bertindak demi kepentingan sempit mereka sendiri dapat menggunakan informasi yang ditransmisikan secara sukarela, bahkan secara sewenang-wenang, untuk merugikan masyarakat. Oleh karena itu, lebih dari sebelumnya, kerja terkoordinasi dari semua kekuatan sosial progresif diperlukan untuk meningkatkan kesadaran warga negara, tidak hanya tentang tingkat perkembangan AI, tetapi juga tentang tren ekonomi, sosial, dan politik dunia kontemporer. Semakin tinggi tingkat kesadaran ini, semakin baik pula kesadaran warga negara akan terlindungi dari ancaman MUI di bidang PS.

3.8 KESIMPULAN

Materi yang dikaji, termasuk yang analisisnya tidak disertakan dalam bab ini, menunjukkan bahwa, baik di tingkat nasional maupun internasional, belum terdapat mekanisme yang komprehensif untuk pengaturan hukum penggunaan AI, termasuk peraturan untuk mencegah tindakan MUI atau untuk menuntut individu dan kelompok yang melakukan tindakan tersebut. Undang-undang nasional beberapa negara memperkenalkan regulasi untuk penggunaan teknologi tertentu, seperti pembuatan deepfake, tetapi, secara umum, aspek hukum pengaturan penggunaan AI dalam sistem PS masih dalam pembahasan.

Kemampuan AI untuk membuat keputusan otonom menimbulkan pertanyaan tentang subjektivitas AI bagi pembuat undang-undang (untuk menentukan siapa yang bertanggung jawab atas pelanggaran: AI itu sendiri, pengembangnya, atau pengguna yang telah melatih AI menggunakan data yang terdistorsi), dan pertanyaan ini hanya dapat diselesaikan dengan kerja sama yang erat antara spesialis AI, pakar PS, psikolog, dan pengacara di tingkat nasional. Upaya pertama sebenarnya sedang dilakukan untuk mengatur penggunaan AI. Dalam kondisi ini, sangat sulit bagi lembaga internasional untuk mengembangkan langkah-langkah tersebut hingga negara-negara dengan indikator perkembangan AI tertinggi membuat norma mereka sendiri dan mengajukan proposal tentang isu hukum internasional berdasarkan norma tersebut. Sebagai contoh, penting untuk dicatat bahwa di tingkat PBB pada saat bab ini ditulis, belum ada resolusi dengan instruksi untuk mengatur deepfake. Semakin penting untuk memulai studi komprehensif tentang MUAI hari ini sebelum MUAI secara komprehensif merusak dan mengacaukan masyarakat. Pandemi COVID-19 telah menjadi katalisator yang kuat bagi digitalisasi berbagai layanan, yang mau tidak mau mengarah pada pertumbuhan MUAI (hal ini telah dibuktikan, misalnya, dengan meningkatnya jumlah serangan phishing). Dalam keadaan ini, para pembuat undang-undang terpaksa bekerja dalam konteks krisis ganda – pandemi dan meningkatnya ketegangan geopolitik.

Dalam lingkungan baru ini, penelitian komprehensif tentang MUAI dalam konteks hak asasi manusia sangat penting bagi sistem PS, dengan kombinasi hasil kerja di dua bidang: (1) arah 'teknis' – analisis skenario situasi di mana MUAI terjadi, berdasarkan kemampuan teknis AI itu sendiri; (2) arah 'sosial' – analisis perkembangan sosial yang memperhitungkan kemungkinan ancaman terhadap hak-hak sipil dan kebebasan dan, dalam konteks ini, risiko spesifik penggunaan AI.

Sintesis hasil kerja di bidang-bidang ini akan memungkinkan kita untuk merumuskan skenario yang menguntungkan dan merugikan bagi implementasi AI dalam konteks pembangunan berkelanjutan.

Dalam mengembangkan pendekatan yang berorientasi pada manusia terhadap permasalahan AI, penting untuk merumuskan tujuan pengembangan teknologi dengan tepat. Penetapan tujuan dan sasaran yang berorientasi sosial akan lebih efektif dengan menggunakan indikator-indikator terencana yang spesifik (seperti Tujuan Pembangunan Berkelanjutan PBB). Disarankan untuk menyempurnakan formulasi kebijakan umum sesuai dengan tujuan sosial. Pendekatan yang berorientasi pada sosial menghindari objektifikasi warga negara dan menjaga subjektivitas mereka dalam mengatur penggunaan AI dalam konteks pembangunan berkelanjutan.

BAB 4

PEMANFAATAN AI UNTUK KEWARGANEGARAAN MELALUI INVESTASI

4.1 PENDAHULUAN

Kewarganegaraan melalui investasi (CBI) adalah proses di mana suatu negara memberikan kewarganegaraan kepada pemohon dengan imbalan investasi dalam perekonomian negara tersebut. CBI berbeda dari skema kependudukan lainnya karena tidak hanya memberikan pemohon status kependudukan dan hak untuk bekerja, tetapi juga paspor dengan masa tinggal tak terbatas yang mencakup hak politik dan kewajiban tradisional di wilayah nasional masing-masing.

Jenis investasi yang diwajibkan dan jumlah uang yang terlibat bervariasi di berbagai negara tuan rumah CBI, meskipun sebagian besar negara yang menawarkan CBI, seperti Antigua dan Barbuda atau Malta, memberi pemohon pilihan antara memberikan kontribusi kepada dana pembangunan nasional yang dikelola pemerintah atau pembelian real estat yang signifikan atau investasi perusahaan.

Program CBI bersifat unik karena tidak sesuai dengan model naturalisasi dan kewarganegaraan tradisional. CBI merupakan bentuk naturalisasi yang dipercepat melalui beberapa cara investasi, dan penggunaannya telah dianggap oleh banyak pelaku industri sebagai akses unik ke kewarganegaraan kedua, alih-alih melalui perolehan hak kelahiran, baik melalui garis keturunan (*jus sanguinis*) maupun kelahiran di wilayah tersebut (*jus soli*).

Alasan calon pemohon memperoleh kewarganegaraan kedua dapat bergantung pada banyak faktor, seperti akses bebas visa ke negara lain, pendidikan dan layanan kesehatan bagi anak-anak pemohon, atau manfaat ekonomi, termasuk kebebasan pergerakan modal tanpa batasan. Mungkin juga terdapat faktor pendorong yang memengaruhi pemohon untuk memperoleh CBI. Misalnya, ketika suatu negara mulai menghadapi keadaan darurat, krisis ekonomi, dan/atau gejolak politik, seringkali orang kaya adalah yang pertama bermigrasi ke luar negeri untuk mencari lingkungan yang lebih aman.

Istilah 'individu berkekayaan tinggi' digunakan oleh beberapa segmen industri jasa keuangan untuk menggambarkan seseorang yang kekayaannya (aset seperti saham dan obligasi) secara signifikan melebihi jumlah tertentu. Bagi mereka, kewarganegaraan kedua merupakan pilihan investasi yang menarik. Program semacam itu telah ada selama beberapa dekade; misalnya, Amerika Serikat, Inggris Raya, dan Kanada telah memiliki 'Program Investor Imigran' sejak pertengahan 1980-an atau pertengahan 1990-an. Negara-negara kecil telah menawarkan jalur yang lebih langsung menuju kewarganegaraan tanpa, atau berdasarkan persyaratan tempat tinggal yang sangat terbatas. Saat ini, negara-negara tersebut mencakup negara-negara seperti Siprus, Dominika, dan St. Kitts dan Nevis; namun, di masa lalu, program tersebut mencakup Irlandia dan beberapa Kepulauan Pasifik. Dalam hal ini, program CBI dapat sangat berbeda dalam cakupan dan karakternya, karena setiap negara telah membentuk

programnya sesuai dengan kebutuhan dan prioritas spesifiknya. Semua program ini bertujuan untuk merangsang pertumbuhan dan lapangan kerja dengan menarik lebih banyak modal dan investasi asing dengan menawarkan kewarganegaraan yang dipercepat kepada individu berpenghasilan tinggi.

Namun demikian, fasilitasi program CBI berpotensi menciptakan risiko kepatuhan perusahaan dan imigrasi yang serius bagi pemerintah, termasuk potensi: (a) pencucian uang, dan (b) risiko keamanan perbatasan.

Dengan mempertimbangkan hal tersebut di atas, uji tuntas perusahaan dan imigrasi merupakan area yang saat ini sedang ditransformasikan oleh AI, tetapi belum diperkenalkan secara sepihak ke dalam program CBI. Uji tuntas perusahaan dan imigrasi yang seragam, atau ketiadaannya, merupakan area risiko utama dalam program CBI, khususnya terkait dengan fasilitasi risiko pencucian uang dan keamanan perbatasan. Sebelum mempertimbangkan uji tuntas terpadu yang didukung oleh AI, penting untuk mengilustrasikan justifikasi dan pentingnya uji tuntas dalam program CBI.

4.2 UJI TUNTAS

Dalam konteks program CBI, uji tuntas dapat digambarkan sebagai investigasi, audit, atau peninjauan terhadap pemohon CBI untuk mengonfirmasi keabsahan permohonan dan menilai potensi risiko yang sedang dipertimbangkan. Pada tingkat paling dasarnya, uji tuntas adalah serangkaian proses investigasi dan audit yang mengidentifikasi, menguatkan, dan memverifikasi informasi tentang individu atau badan usaha sehingga negara tuan rumah CBI dapat dilindungi dari pencucian uang dan risiko keamanan perbatasan. Uji tuntas CBI harus mengkhususkan diri dalam pemeriksaan latar belakang menyeluruh pada catatan kriminal, korupsi, sanksi, litigasi, aset yang dibekukan, penyembunyian, penipuan pajak, kejahatan keuangan dan penggelapan, melalui basis data polisi, dll. Terlepas dari risiko-risiko ini, pembenaran untuk memperkuat uji tuntas dalam kepatuhan perusahaan dan imigrasi dengan alat AI terutama dimotivasi oleh jumlah aplikasi, yang menunjukkan pertumbuhan industri yang cepat, dan pendapatan signifikan yang dihasilkan oleh pemerintah CBI. Potensi risiko diperkuat melalui meningkatnya jumlah pelamar dan potensi satu atau lebih pelamar yang jahat untuk menghindari metode uji tuntas saat ini dan merusak reputasi program.

Secara statistik, di wilayah Karibia, aplikasi telah meningkat dalam banyak kasus hingga tahun 2020. Ini dicontohkan dalam Pernyataan Anggaran 2021 untuk Antigua dan Barbuda, yang disampaikan oleh Perdana Menteri Yang Terhormat Gaston Browne. Dalam pidatonya, ia menguraikan, antara lain, pentingnya pendapatan dari program kewarganegaraan mereka. Pada tahun 2020, total 366 aplikasi diterima, mewakili penurunan 22% dari tahun sebelumnya. Meskipun demikian, CBI menghasilkan Rp. 1.157 Triliun juta pada tahun 2020 saja. Sejak awal berdirinya pada tahun 2012, Antigua dan Barbuda telah menaturalisasi lebih dari 4000 warga negara melalui program mereka. Hanya 188 km ke selatan, di Dominika, unit kewarganegaraan berdasarkan investasi (CIU) mereka melaporkan 2.100 aplikasi ke CBI antara Juli 2018 dan Juni 2019. Jumlah ini mewakili angka persetujuan aplikasi tertinggi yang tercatat secara resmi dalam satu tahun di antara CBI Karibia. Dalam premis ini, perlu dicatat bahwa St. Kitts dan

Nevis, yang pemerintahnya menolak untuk mengungkapkan data tahunan yang tepat karena dianggap sebagai rahasia dagang, mungkin telah menyaingi Dominika pada tahun 2019. Terlepas dari itu, Dominika menerbitkan 5.815 paspor untuk investor dan anggota keluarga mereka antara tahun 2017 dan 2020 dan menghasilkan Rp. 12 Triliun bagi negara tersebut. Sebagai kesimpulan, di St. Lucia, menurut laporan berita, lebih dari 1.000 orang telah diberikan kewarganegaraan St. Lucia pada tahun 2020, tahun kelima operasinya, dan ini telah mengumpulkan Rp. 11 Triliun.

Di Uni Eropa, Bulgaria, Siprus, dan Malta adalah satu-satunya Negara Anggota yang mengoperasikan skema CBI. Malta dapat dimengerti menetapkan batasan tahunan pada jumlah kewarganegaraan oleh pemohon investasi yang dapat dinaturalisasi. CBI mereka menaturalisasi total 3.673 pelamar antara awal berdirinya pada tahun 2013 dan 2018. Malta menghasilkan lebih dari Rp. 900 Triliun sejak dimulainya CBI, yang mewakili 7,2% dari Produk Domestik Bruto Malta tahun 2017. Sebagai perbandingan, Kementerian Keuangan Siprus menyatakan hasil CBI mereka sejak awal berdirinya, dan menurut hasil audit antara Juni 2013 dan Desember 2019, 2.855 investor menerima paspor Siprus untuk investasi, dan dalam 7 tahun program tersebut menghasilkan Rp. 97 Triliun.

Secara global, CBI telah tumbuh menjadi industri yang diperkirakan bernilai Rp. 3 Triliun, dengan, dilaporkan, 9.000 pemohon CBI disetujui di seluruh dunia pada tahun 2018/2019. Diperkirakan bahwa pada tahun 2020 di Turki saja terdapat sekitar 13.000–14.000 pemohon, meskipun terjadi pandemi COVID-19. Angka ini diperkirakan akan bertambah karena program-program menjadi jauh lebih terjangkau dan sebagai hasil dari persaingan CBI untuk menarik lebih banyak individu kaya. Penurunan harga ini juga menimbulkan masalah karena dapat memberikan akses yang lebih mudah bagi individu dengan niat jahat. Disarankan agar status quo uji tuntas saat ini dalam industri CBI didukung dengan berbagai alat pembelajaran mesin yang meningkatkan efektivitas dan efisiensinya dengan mempertimbangkan pertumbuhan industri yang diilustrasikan di atas, dan risiko yang diperkuat yang dihadirkan oleh pertumbuhan yang diperkirakan ini.

Untuk mendukung peningkatan uji tuntas di CBI, Investment Migration Council, berkoordinasi dengan BDO USA (layanan jaminan, pajak, dan konsultasi), Exiger (perusahaan regulasi global dan kejahatan keuangan, risiko, dan kepatuhan), dan Refinitiv (lembaga data dan infrastruktur pasar keuangan), membentuk kelompok kerja uji tuntas untuk memeriksa status uji tuntas di seluruh sektor migrasi investasi dan mengeksplorasi potensi standar minimal, transparansi yang lebih besar, dan berbagi informasi di seluruh industri. Kelompok kerja ini melakukan dua laporan, yang pertama memberikan tinjauan kritis tentang proses uji tuntas dan penilaian apakah praktik uji tuntas saat ini perlu ditingkatkan. Setelah itu, laporan kedua menggambarkan rekomendasi untuk standar minimum dasar bagi penyedia uji tuntas dalam industri. Menjadi jelas dalam rekomendasi ini bahwa pengenalan alat AI tertentu menjadi lebih realistis untuk berpotensi menetapkan standar minimum dalam industri.

Saat ini, ada sedikit atau tidak ada koordinasi uji tuntas atau standar yang seragam antara negara-negara aktif CBI dan AI berpotensi membantu dalam hal ini. Berdasarkan premis ini, saat ini, agen berlisensi diwajibkan untuk melakukan administrasi awal pemeriksaan uji

tuntas dan menyerahkan aplikasi mereka kepada CIU pemerintah. Dengan demikian, agen memiliki kesempatan pertama untuk mengidentifikasi dan menolak pelamar yang tidak memenuhi persyaratan uji tuntas, dan apakah agen membuat keputusan awal yang memadai sebagian bergantung pada kualitas uji tuntas yang mereka lakukan.

CIU harus, antara lain, melakukan verifikasi validitas catatan kepolisian dan pengadilan individu serta memeriksa basis data nasional dan catatan pemerintah lainnya. Teknik pembelajaran mesin dan algoritma yang digunakan dalam perangkat lunak seperti 'Veriff' dapat membantu dalam hal ini. Selain itu, CIU harus berkoordinasi dengan lembaga penegak hukum internasional untuk mendapatkan detail tentang surat perintah yang beredar atau kecurigaan aktivitas kejahatan internasional. CIU harus berfokus pada integritas dan keamanan secara keseluruhan, tidak hanya program individual, tetapi juga perspektif industri secara keseluruhan. Disarankan agar CIU menggabungkan perangkat yang paling inovatif dan canggih secara teknis untuk melakukan uji tuntas bertingkat dan terkoordinasi secara multilateral.

Tugas uji tuntas dapat menimbulkan beban administratif yang signifikan bagi agen berlisensi dan CIU, karena mereka dituntut untuk tidak hanya teliti tetapi juga proaktif, mengingat saat ini di CBI belum ada kolaborasi terkait uji tuntas korporat dan imigrasi.

Implementasi AI kemungkinan akan menawarkan peluang besar di bidang ini. Perangkat pembelajaran mesin akan membantu meringankan beban kerja manual CIU dan dapat membantu agen dalam melakukan uji tuntas awal. Selain itu, teknik pengenalan wajah otomatis dalam konteks biometrik imigrasi dan keamanan perbatasan dapat secara efisien memaksimalkan proses uji tuntas manual yang saat ini dilakukan. Penyaringan biometrik pelamar oleh CIU sangat penting untuk menjaga keamanan dan menjaga kredibilitas masing-masing program CBI.

4.3 KEPATUHAN KORPORASI DALAM CBI

Investasi asing langsung yang dihasilkan oleh program-program CBI ini mungkin substansial; namun, hal ini juga dapat menimbulkan implikasi makroekonomi yang signifikan di berbagai sektor ekonomi negara CBI. Misalnya, arus masuk di St Kitts dan Nevis, dan pada tingkat yang lebih rendah di Dominika, telah tumbuh menjadi bagian yang signifikan dari Produk Domestik Bruto, yang memengaruhi permintaan agregat dan menimbulkan pertanyaan tentang risiko terhadap stabilitas makroekonomi dan keberlanjutan fiskal. Tergantung pada programnya, ada tiga jenis arus masuk utama:

- i) kontribusi untuk biaya aplikasi terkait pemerintah;
- ii) kontribusi yang tidak dapat dikembalikan kepada pemerintah atau dana kuasi-pemerintah (misalnya Dana Pembangunan Nasional); dan
- iii) investasi di sektor swasta atau publik, yang seringkali dapat dijual atau ditebus setelah jangka waktu tertentu.

Di luar potensi risiko makroekonomi, CBI, jika dieksploitasi, akan membantu menyembunyikan atau memfasilitasi kejahatan keuangan, ekonomi, dan/atau terorganisir, yang meliputi penyuapan, korupsi, dan/atau pencucian uang. Masalah terakhir ini signifikan. Pencucian Uang dapat didefinisikan sebagai 'proses yang digunakan penjahat untuk menyembunyikan,

mengendalikan, menginvestasikan, dan mendapatkan keuntungan dari hasil kegiatan kriminal mereka. Hal ini terjadi ketika penjahat mencoba menciptakan latar belakang yang sah untuk uang mereka'. Algoritma pembelajaran mesin telah terbukti sangat berguna dalam memantau transaksi dan aktivitas mencurigakan. Kedua aspek kunci ini telah diakui oleh Gugus Tugas Aksi Keuangan dan mayoritas Negara Anggota OECD, sebagai karakteristik utama pencucian uang, dan dengan demikian merupakan inti dari standar anti pencucian uang nasional.

Hingga baru-baru ini, lembaga korporat mengandalkan sistem pemantauan transaksi anti pencucian uang tradisional berbasis aturan dan penyaringan nama, yang menghasilkan angka positif palsu yang tinggi karena ambang batas aturan. Singapura, sebagai pusat keuangan global terkemuka, memiliki 'kursi baris depan' untuk ancaman pencucian uang ini. Faktanya, negara ini telah memimpin dalam mengatasi 'lanskap ancaman' yang berkembang melalui inisiatif, solusi, dan forum yang inovatif, seperti yang terlihat dalam penyelenggaraan Singapore Fintech Festival yang berkelanjutan oleh Otoritas Moneter Singapura. Inovasi dalam kepatuhan CBI diperlukan baik untuk mengurangi positif palsu dalam aplikasi maupun untuk menghasilkan efektivitas yang lebih besar dalam cara risiko pencucian uang dipantau dan ditangani oleh CIU.

Dalam CBI, teknik pembelajaran mesin seperti deteksi anomali dapat digunakan untuk mengidentifikasi pola transaksi yang sebelumnya tidak terdeteksi, anomali data, dan hubungan di antara individu dan entitas yang mencurigakan. Program berbasis AI seperti Luminance dapat diajarkan untuk mengenali perilaku mencurigakan dan risiko investasi, dan memeringkatnya sesuai dengan itu. Tantangan umum dalam program CBI adalah pemantauan transaksi dan peringatan aktivitas mencurigakan, dan di sinilah pembelajaran mesin dapat mengajarkan komputer untuk mengklasifikasikan peringatan uji tuntas sebagai risiko tinggi, sedang, atau rendah.

Investasi perusahaan dan fasilitasi transaksi internasional merupakan aktivitas bisnis utama yang berdampak signifikan terhadap perekonomian nasional. Aplikasi AI yang digunakan untuk uji tuntas juga dapat membantu membuat keputusan yang tepat tentang pelamar CBI dan dapat dilatih untuk mengidentifikasi karakteristik perilaku atau indikator yang menyoroti ketika aktivitas benar-benar mencurigakan. Pembelajaran mesin dapat diterapkan pada penyaringan nama, di mana agen dan CIU diwajibkan untuk menyaring pelamar berdasarkan daftar global penjahat yang diketahui dan organisasi yang masuk daftar hitam atau dikenai sanksi. Dalam status quo saat ini, pendekatan berbasis aturan bersifat memberatkan dan manual, yang mengakibatkan peningkatan beban kerja untuk kepatuhan, serta potensi kesenjangan dalam pengawasan dan pemantauan.

Beberapa area lain yang mendapatkan daya tarik termasuk deteksi penipuan, pelaporan otomatis, dan pengawasan yang ditingkatkan, termasuk pemantauan transaksi berbasis suara, video, teks, dan pola. Teknologi pembelajaran mesin dalam deteksi penipuan saat ini memungkinkan bank untuk memprediksi secara akurat apakah suatu rekening berisiko dibobol, dan diusulkan agar teknik serupa diadaptasi dalam program CBI untuk memprediksi secara akurat apakah pemohon CBI memiliki risiko keuangan. Perusahaan seperti 'DataVisor' menyediakan perangkat lunak khusus yang dapat melacak transaksi penipuan. DataVisor

menggabungkan kemampuan pembelajaran mesin terapan dengan alur kerja investigasi yang canggih dan jaringan intelijen lebih dari 4 miliar akun pengguna untuk memberikan sinyal penipuan waktu nyata, wawasan, dan perlindungan guna menjaga kepercayaan dan keamanan yang vital. Pada tahun 2020, perusahaan tersebut mengklaim bahwa perangkat lunak tersebut mendeteksi hingga 30% lebih banyak penipuan dengan akurasi hingga 90% lebih tinggi.

Pemerintah CBI biasanya menghabiskan sumber daya yang besar untuk informasi tentang calon pemohon, terutama dalam konteks investasi program dan persyaratan administrasi uji tuntas. Pengenalan chatbot AI ke CBI berpotensi meminimalkan pengeluaran dan meningkatkan efisiensi dalam hal ini, sehingga membantu meningkatkan pendapatan program. Chatbot adalah asisten daring virtual waktu nyata yang dapat mensimulasikan perilaku percakapan manusia. Dalam 'percakapan' daring dengan calon pelamar, chatbot dapat bertanya dan menjawab pertanyaan, dan umumnya berinteraksi secara 'cerdas'. Data dari obrolan daring ini dapat dikumpulkan dan dianalisis untuk menghasilkan hasil yang lebih baik di masa mendatang. Kesimpulannya, semua sistem pemantauan dan analitik, bukan hanya aplikasi pembelajaran mesin, bergantung pada data berkualitas tinggi. Mengingat ketersediaan data perusahaan yang belum pernah terjadi sebelumnya secara global, dapat dikatakan terdapat minat baru yang meluas untuk menerapkan metode pembelajaran mesin berbasis data pada permasalahan yang pengembangan solusi perusahaan konvensional masih terkendala.

4.4 KEPATUHAN IMIGRASI DALAM CBI

'Manajemen Perbatasan' berkaitan dengan administrasi kebijakan imigrasi. Meskipun makna tepatnya dapat bervariasi menurut program CBI nasional, biasanya berkaitan dengan aturan, teknik, dan prosedur yang mengatur aktivitas dan lalu lintas wilayah atau zona perbatasan yang telah ditentukan. Kepatuhan Imigrasi secara bertahap digunakan untuk melakukan tugas-tugas tertentu dalam hal ini, termasuk pemeriksaan identitas, keamanan dan kontrol perbatasan, serta analisis data tentang pemohon visa dan suaka. Terdapat contoh penggunaan Kepatuhan Imigrasi semacam itu di Kanada dan Jerman. Dalam konteks program CBI, risiko keamanan unik tertentu diidentifikasi, termasuk pemohon yang melarikan diri dari keadilan di yurisdiksi asal mereka, dan menyalahgunakan hak istimewa perjalanan bebas visa yang diberikan melalui paspor yang dibeli. Hilangnya hak istimewa perjalanan bebas visa dicontohkan dalam pencabutan akses bebas visa tahun 2014 antara St Kitts dan Nevis dan Kanada untuk semua warga negara St Kitts dan Nevis karena praktik uji tuntas, atau ketiadaannya, dalam program CBI mereka. Dalam hal ini, dikemukakan bahwa satu pemohon yang melakukan kejahatan berpotensi menempatkan risiko dan kewajiban yang tak terbantahkan tidak hanya pada program domestik tetapi juga pada CBI secara umum.

Contoh menarik dari pelarian keadilan diajukan oleh kasus pengusaha India Mehul Choski tahun 2019, yang dituduh di India atas dugaan kegiatan penipuan terhadap Punjab National Bank. Meskipun ia memegang kewarganegaraan Antigua dan Barbuda melalui CBI mereka sejak 2017, dan ada perjanjian ekstradisi antara kedua negara ini, Tn. Choski memulai proses hukum di Antigua dan Barbuda terhadap ekstradisinya ke India dan menyerahkan

kewarganegaraan India-nya. Sebagaimana dilaporkan pada tahun 2021, Perdana Menteri Antigua dan Barbuda, Gaston Browne, tetap bersikeras bahwa Meहुल Choksi harus meninggalkan negara tersebut. Hal ini karena ia merusak citra program CBI negara tersebut. Saat artikel ini ditulis, kasus ini masih dalam proses gugatan hukum oleh Bapak Choksi di pengadilan Antigua dan Barbuda.

Dalam yurisdiksi praktik terbaik, uji tuntas pelancong diinformasikan oleh pembaca dokumen, yang menyediakan mekanisme yang efisien dan akurat untuk mengekstrak data dari dokumen perjalanan, secara otomatis memicu pencarian daftar pantauan, memungkinkan verifikasi identitas biografis dan biometrik, dan merekam masuknya pelancong. Dalam ekspresi akhir dari otomatisasi ini, pelancong berinteraksi dengan 'eGates' swalayan dan kios saat masuk (atau keberangkatan) tanpa memproses masukan dari staf badan perbatasan, sehingga melepaskan sumber daya untuk penempatan kembali ke tujuan keamanan atau fasilitas lainnya. Saat ini, '*proyek AI e-Government*' semacam itu ada di Kanada dan Jerman. Mereka menyediakan berbagai fungsi AI untuk mendukung penilaian otomatis berbagai jenis aplikasi yang diajukan untuk tujuan kepatuhan imigrasi.

Di Kanada, pemerintah telah dalam proses mengembangkan sistem 'analisis prediktif' untuk mengotomatiskan kegiatan uji tuntas yang saat ini dilakukan oleh pejabat imigrasi, dan untuk mendukung evaluasi aplikasi imigran dan pengunjung. Pemerintah juga telah meminta masukan dari sektor swasta terkait proyek tahun 2018 untuk 'Solusi AI' dalam pengambilan keputusan dan penilaian imigrasi.

Jerman telah melakukan uji coba proyek yang menggunakan teknologi pengenalan wajah dan dialek untuk pengambilan keputusan dalam proses penentuan suaka. Sistem ini menerapkan sejumlah alat, seperti transkripsi nama otomatis, identifikasi dialek otomatis untuk memvalidasi negara asal yang terdaftar, pemeriksaan gambar otomatis untuk mencegah penggandaan berkas pencari suaka, dan pengambilan informasi dari ponsel pintar untuk mengumpulkan informasi identitas lebih lanjut jika dokumen identitas hilang.

Modul AI tersebut dapat diterapkan pada program CBI untuk memfasilitasi uji tuntas bertingkat, dengan standar minimum yang diakui di seluruh industri. Dengan memperkuat uji tuntas dengan AI, agen dan CIU tidak hanya akan menjadi lebih efisien, tetapi juga standar praktik terbaik untuk mitigasi risiko yang disebutkan di atas dapat diwujudkan secara lebih efektif. Layanan ini dapat diintegrasikan ke dalam formulir aplikasi dan sistem pemrosesan CIU yang spesifik, yang mencakup pemrosesan alur kerja, yaitu tugas administratif, dan manajemen dokumen. Layanan yang telah terbukti seperti yang diilustrasikan di atas meliputi penilaian berbasis aturan, saran berbasis skema, penggalan data, penalaran berbasis kasus, dan pembelajaran mesin. Jika diterapkan dalam CBI, layanan ini akan menyederhanakan proses dan alur kerja sekaligus memastikan semua aplikasi diproses secara adil dan akurat, dengan mematuhi semua hukum dan peraturan yang relevan.

4.5 CHALLENGE DENGAN PENGENALAN AI

Setelah implementasi AI pada program CBI, pemerintah perlu merespons potensi tantangan sosial dan hukum yang ditimbulkan oleh AI. Tantangan ini dapat mencakup bias

algoritmik, eksploitasi data, dan ketidaksetaraan terkait siapa yang diberikan kewarganegaraan.

Pada tahun 2020, raksasa internet Amazon menggunakan AI untuk membangun alat penyaringan resume dengan harapan dapat membuat proses evaluasi lamaran lebih efisien. Mereka membangun algoritma menggunakan resume yang telah dikumpulkan perusahaan selama satu dekade. Karena resume tersebut sebagian besar berasal dari pria, sistem Amazon belajar sendiri bahwa kandidat pria lebih disukai dan kemudian perusahaan menghapus alat tersebut. Ini adalah contoh dari apa yang disebut 'bias algoritmik'. Hal ini menimbulkan tantangan serius bagi penggunaan dan tata kelola AI. Hal ini karena pembelajaran mesin dalam program CBI akan melibatkan 'pelatihan' AI, menggunakan contoh dan data tentang bagaimana manusia sebelumnya membuat keputusan administratif. Contoh atau kumpulan data ini mungkin mengandung bias sosial dan, akibatnya, hal ini dapat tercermin dalam keputusan sistem berbasis AI.

Secara hipotetis, jika data yang digunakan untuk pelatihan AI dalam program CBI hanya berisi pelamar dari satu latar belakang etnis dan mencakup risiko spesifik terkait pelamar dari latar belakang tersebut, keputusan AI secara keseluruhan akan mencerminkan bias, sikap, dan opini dalam pertimbangan aplikasi tersebut. Hal yang sama berlaku untuk keputusan lain seperti administrasi investasi perusahaan dan pertimbangan imigrasi.

Penulis menggambarkan bahwa sangat sulit untuk menciptakan pelatihan AI yang netral dan independen dari semua bias manusia. Salah satu pendekatan untuk netralitas bias adalah dengan menggunakan contoh data keputusan dari manusia dengan latar belakang budaya dan sosial yang luas. Ada peningkatan penekanan pada hak asasi manusia dalam percakapan tentang isu-isu etika dan sosial yang diangkat oleh teknologi AI. Dalam pertimbangan CBI, diskriminasi gender merupakan pertimbangan utama.

Gender, dalam konteks negara-negara CBI, mengacu pada peran, perilaku, aktivitas, atribut, dan norma yang dianggap pantas bagi pria dan wanita oleh suatu masyarakat pada waktu tertentu. Atribut, peluang, dan hubungan ini dikonstruksi secara sosial dan dipelajari melalui proses sosialisasi. Tujuannya adalah untuk mengidentifikasi faktor-faktor gender yang memengaruhi perilaku investasi dan imigrasi serta untuk menentukan proses pengambilan keputusan investasi, khususnya yang berkaitan dengan pelamar perempuan. Penelitian empiris tentang perbedaan gender dalam pengambilan risiko memang menunjukkan bahwa perempuan lebih sedikit mengambil risiko dibandingkan laki-laki.

Terkait hal ini, disarankan agar program CBI mempertimbangkan risiko kepatuhan investasi dan imigrasi yang terkait secara khusus dengan perempuan berpenghasilan tinggi dan kontrol perbatasan kedaulatan. Integrasi kebijakan gender ke dalam kontrol perbatasan AI dalam CBI sangat penting untuk mematuhi kerangka hukum internasional dan regional, instrumen dan norma terkait keamanan, kesetaraan gender, dan hak asasi manusia. Akibatnya, sistem AI, tidak terbatas pada CBI, harus dirancang dan dioperasikan agar sesuai dengan prinsip-prinsip martabat manusia, hak dan kebebasan fundamental, serta keberagaman budaya, sehingga menghindari bias algoritmik. Pengenalan AI menghadirkan tantangan dan risiko sosial baru yang memerlukan respons terkoordinasi.

Mengingat karakternya yang multidimensi, AI secara inheren menyentuh spektrum penuh bidang hukum, mulai dari filsafat hukum, hak asasi manusia, hukum kontrak, hukum perdata, hukum ketenagakerjaan, hukum pidana, hukum pajak, hukum investasi, dan hukum acara, untuk menyebutkan beberapa di antaranya. Selain itu, pengumpulan data juga menimbulkan masalah yurisdiksi. Mustahil untuk membahas semua aspek ini secara rinci dalam satu bab. Beberapa isu ini dibahas dalam bab-bab lain dalam volume ini.

4.6 KESIMPULAN

Sistem AI yang saat ini digunakan dalam kepatuhan korporat dan imigrasi dapat diterapkan pada program CBI. Dalam kepatuhan korporat, AI dapat membantu mengidentifikasi pelamar yang melakukan pelanggaran, dan mereka yang terlibat dengan hasil kejahatan, termasuk pencucian uang. Dengan demikian, uji tuntas korporat di CIU akan terbukti lebih efektif dan efisien. Sebagai alternatif, untuk imigrasi, tugas administratif biometrik sederhana dan uji tuntas keamanan perbatasan secara keseluruhan akan jauh lebih baik, sejalan dengan modul AI yang telah digunakan di sektor imigrasi di Kanada dan Jerman. Dalam hal ini, data harus diunggah tanpa bias algoritmik dan mewakili dinamika sebenarnya dari pertimbangan spesifik program CBI terkait bias gender.

Penting untuk merumuskan solusi untuk penilaian dan mitigasi risiko sambil menjaga keseimbangan antara non-diskriminasi, keamanan perbatasan, dan peluang investasi. CIU harus, antara lain, berfokus pada pencucian uang dan keamanan perbatasan untuk memastikan bahwa program mereka memberikan keuntungan jangka panjang yang berkelanjutan dengan menggunakan perangkat uji tuntas yang paling inovatif dan berteknologi canggih. Pemerintah nasional yang mempraktikkan CBI harus secara keseluruhan menghormati standar hukum global yang ditetapkan oleh lembaga-lembaga seperti Komisi Eropa, serta standar hukum internasional yang relevan, sebagaimana disebutkan sebelumnya, yang dirancang untuk mencegah pencucian uang dan risiko lainnya.

BAB 5

POTENSI, TANTANGAN, REGULASI AI BIDANG KESEHATAN

Pengambilan keputusan perawatan kesehatan merupakan salah satu area yang terdampak oleh AI, di mana AI terus terintegrasi ke dalam layanan kesehatan untuk memfasilitasi pemberian perawatan pasien yang lebih baik dan lebih efisien. AI dalam perawatan akhir hayat baru-baru ini menarik perhatian ketika sebuah studi terbaru menyelidiki penerapan pengambilan keputusan berbantuan AI untuk memprediksi peluang pasien untuk bertahan hidup dan pulih. Studi tersebut menyoroti potensi penerapan AI dalam meningkatkan kepercayaan diri dokter dalam penilaian klinis dan kemungkinan bahaya bagi pasien yang timbul dari risiko penggunaan AI. Pengambilan keputusan lanjutan merupakan bagian penting dari perawatan kesehatan untuk memandu keputusan perawatan dan khususnya berguna untuk merencanakan perawatan dan pengobatan di masa mendatang, mulai dari penyakit jangka panjang, seperti penyakit progresif, hingga penyakit terminal, seperti kanker. AI sering digunakan dalam perencanaan perawatan akhir hayat, sehingga penting untuk mengkaji penerapan AI di area ini.

5.1 AI DALAM LAYANAN KESEHATAN

AI semakin banyak diterapkan dalam penyediaan layanan kesehatan di berbagai disiplin ilmu, seperti kanker, neurologi, kardiologi, urologi, sistem pasien, diagnostik, dan pencitraan medis. AI sering digunakan untuk mengelola dan menganalisis data, membuat keputusan, menyalin informasi, dan membantu komunikasi dokter-pasien. Contoh aplikasi spesifik di seluruh sektor layanan kesehatan meliputi interpretasi citra dalam pencitraan medis, prediksi, diagnosis, dan prognosis stroke dini, diagnosis syok septik dan pengobatan penyakit paru obstruktif kronik, pengambilan keputusan di akhir hayat, layanan kesehatan personal dan diagnosis tambahan (seperti subkelompok diabetes atau autisme), dan dalam ranah layanan kesehatan primer dengan pemodelan prediktif pada data kesehatan dan analitik bisnis untuk penyedia layanan kesehatan primer. Penyedia layanan kesehatan swasta telah menguji coba AI di rumah sakit NHS untuk menganalisis hasil tes, menentukan perawatan yang tepat, dan memastikan peningkatan perawatan pasien ke spesialis yang tepat dengan segera.

Selain berharga untuk mengidentifikasi pola perawatan dan diagnosis melalui penggalian data rekam medis digital, AI digunakan untuk mendukung tenaga kesehatan profesional dalam membuat keputusan yang lebih baik melalui sistem pendukung keputusan klinis. Sistem ini menganalisis data pasien yang relevan untuk menyarankan area yang perlu diperhatikan atau potensi komplikasi dan mengusulkan jalur perawatan, mulai dari praktik perawatan pasien pascaoperasi, pengobatan (kisaran, dosis, alergi, dan interaksi dengan resep lain), pemantauan, dan tindak lanjut berkala dalam lingkungan yang hemat biaya dan aman. Aplikasi ini menawarkan manfaat 'mengurangi kesalahan medis [dan] meningkatkan konsistensi dan efisiensi layanan kesehatan', sehingga mengurangi beban dokter dalam mendiagnosis kasus-kasus kompleks. Penerapannya dalam respons ancaman kesehatan

masyarakat telah diketahui, seperti dalam epidemi Ebola, untuk mengidentifikasi laporan penyakit pasien, yang dicocokkan dengan riwayat perjalanan mereka, melalui pemrosesan cepat informasi yang luas terkait penyakit dan perawatan.

AI juga digunakan dalam pengobatan yang dipersonalisasi. Contohnya adalah kemitraan antara Auckland University of Technology dan Knowledge Engineering Discovery Research Institute yang telah menghasilkan tingkat keberhasilan tinggi dalam memprediksi stroke pada pasien: '95 persen untuk satu hari ke depan, dan 70 persen untuk 7 dan 11 hari sebelum stroke terjadi'. Hal ini bermanfaat untuk tindakan pencegahan guna mengelola kecacatan jangka panjang atau mencegah kematian, dan menghemat biaya dalam jangka panjang. Akurasi luar biasa seperti itu juga telah muncul di bidang perawatan kesehatan spesifik lainnya, misalnya, presisi hingga 80% dalam deteksi penyakit jantung dan akurasi 94% dalam endoskopi pertumbuhan kanker Jepang. AI memiliki banyak hal untuk ditawarkan dalam perawatan kesehatan di masa-masa biasa, tetapi bagaimana dengan kegunaannya di masa-masa luar biasa, seperti pandemi COVID-19 yang telah melanda dunia pada tahun 2020? Pandemi menciptakan rasa urgensi yang mendalam dalam populasi untuk mengurus masalah perawatan kesehatan dan untuk pengambilan keputusan untuk perawatan medis jika individu terinfeksi COVID-19. Sifat COVID-19, yang terutama memengaruhi fungsi pernapasan, telah memperkuat keputusan klinis yang sulit ketika memutuskan apakah akan melakukan resusitasi pasien, dan ketika menarik atau menahan dukungan ventilasi untuk pasien yang dirawat di rumah sakit. AI dapat memainkan peran penting dalam membantu dokter dan penyedia layanan kesehatan membuat keputusan ini, misalnya, dalam mengidentifikasi dengan cepat lintasan kesehatan dan penyakit berdasarkan profil pasien menggunakan bukti ilmiah terbaru yang tersedia tentang bagaimana COVID-19 berkembang dan memengaruhi individu. Jaringan yang luas dan fitur pembelajaran mendalam AI berguna untuk memproses secara bersamaan beberapa basis data ilmiah yang berkaitan dengan COVID-19 dan pilihan perawatan yang dilakukan di seluruh dunia. Selain menavigasi basis data berbasis bukti ini, AI dapat digunakan untuk mengakses informasi yang lebih luas, seperti perjalanan pasien dan riwayat medis dan kemungkinan bertahan hidup, untuk menyajikan profil yang lebih jelas tentang hasil klinis dan saran kesehatan potensial. Potensi AI tampaknya tak terbatas. Bagian selanjutnya akan mengeksplorasi potensinya dalam mendukung pengambilan keputusan klinis, dengan fokus pada pengambilan keputusan tingkat lanjut di bidang kesehatan.

5.2 POTENSI AI DALAM KEPUTUSAN PERAWATAN KESEHATAN

Pengambilan keputusan di muka merupakan bagian dari penyediaan layanan kesehatan dan sebagian besar digunakan dalam perencanaan perawatan akhir hayat jika terjadi ketidakmampuan. Keputusan di muka memberikan kesempatan bagi seseorang untuk menyetujui atau menolak perawatan di kemudian hari ketika mereka tidak mampu melakukannya. Keputusan yang dibuat di muka dapat dinyatakan secara lisan atau tertulis. Ada dua aspek utama dalam bentuk pengambilan keputusan ini. Aspek pertama berkaitan dengan validitas pada saat keputusan dibuat dan aspek kedua berkaitan dengan penerapannya ketika keputusan tersebut hendak diimplementasikan. Agar dokter dapat melaksanakan keinginan

pasien yang diungkapkan, keputusan tersebut harus valid dan dapat diterapkan. Keputusan di muka dianggap valid jika dibuat ketika orang tersebut memiliki kapasitas mental untuk melakukannya, telah membuat keputusan secara sukarela, dan dengan pemahaman penuh tentang sifat keputusan tersebut. Ketika keputusan tersebut diimplementasikan, akan ditentukan apakah keadaan yang disajikan masih ada atau apakah keputusan tersebut berada dalam lingkup perawatan, dan bahwa orang tersebut tidak menunjukkan perubahan pikiran apa pun. Beberapa kesulitan yang dihadapi penyedia layanan kesehatan dalam kedua aspek tersebut berkisar dari ketidakmampuan untuk menentukan apakah keputusan telah dibuat secara sah atau masih berlaku untuk keadaan saat ini. Kesulitan-kesulitan ini tidak hanya menimbulkan konflik antara keluarga dan dokter, tetapi juga menyebabkan sengketa hukum.

AI berpotensi meringankan beberapa kesulitan dalam bentuk pengambilan keputusan layanan kesehatan ini. Dalam hal perencanaan dan pengambilan keputusan, AI dapat melakukan tugas-tugas yang membantu proses konsultasi dokter-pasien. Aspek pertama pengambilan keputusan mengharuskan dokter untuk yakin akan kapasitas mental pasien, kesukarelaan, dan pemahamannya terhadap keputusan tersebut. Penilaian kapasitas mental biasanya dilakukan oleh panel konsultan, yang terdiri dari psikolog atau psikiater, untuk membentuk pandangan profesional mengenai kondisi mental individu yang bersangkutan. Dalam menilai kapasitas mental, AI dapat berperan memprediksi kondisi mental berdasarkan data perilaku dan respons pasien terhadap pertanyaan, lalu membandingkannya dengan pengamatan pribadi dokter konsultan untuk mencapai keputusan. Hal ini khususnya bermanfaat ketika pasien memiliki kondisi kesehatan mental yang mendasarinya, seperti depresi atau psikosis, di mana kapasitas mental cenderung berfluktuasi. Penilaian profesional ini, dikombinasikan dengan keluaran dari pemrosesan terintegrasi AI dapat menghasilkan profil yang lebih bernuansa dari kondisi mental individu secara keseluruhan pada waktu tertentu. Hal ini juga bermanfaat dalam situasi di mana terdapat perbedaan pendapat medis tentang kondisi mental pasien, seperti yang sering ditunjukkan dalam proses pengadilan. Hakim sering beralih ke bukti para ahli untuk membentuk kesimpulan tentang apakah persyaratan telah terpenuhi, dengan mengacu pada catatan dan laporan dokter. Algoritma AI menyediakan sumber informasi tambahan untuk mendukung atau membantah temuan. Meskipun aspek AI ini belum diterapkan secara khusus di area ini, masuk akal untuk mendalilkan bahwa AI akan digunakan dalam waktu yang tidak terlalu lama untuk membantu hakim dan pengadilan dalam membuat keputusan tentang kondisi mental individu ketika mereka sedang diinterogasi. Kita sudah bisa melihat beberapa aplikasi AI di bidang psikiatri dan psikologi, di mana teknologi AI digunakan untuk menganalisis dan memprediksi ciri, perilaku, respons, dan kondisi mental tertentu dari individu yang menderita berbagai penyakit mental, seperti depresi, keinginan bunuh diri, atau skizofrenia, dan gangguan progresif seperti demensia.

Pasien perlu diinformasikan tentang prognosis kesehatan mereka untuk merumuskan perawatan medis mereka di masa depan. Hal ini tidak hanya membutuhkan pemahaman tentang kondisi mental orang tersebut tetapi juga kondisi kesehatan, diagnosis, dan prognosis jika mereka menderita suatu penyakit. Sangat mungkin bahwa aplikasi yang dibantu AI akan

bermanfaat dalam membuat rencana perawatan di masa depan bagi individu yang menderita penyakit terminal atau progresif, seperti demensia, Parkinson atau Alzheimer. Fungsionalitas pembelajaran mesin dan pembelajaran mendalam memungkinkan referensi silang yang luas dan mendalam dari catatan kesehatan dan basis data untuk membentuk pola dan profil pasien untuk memberikan hasil yang paling mungkin bagi para dokter. Selain itu, kumpulan data dari penyakit yang dikenali ini dapat dijelaskan melalui basis data penyakit dan dicocokkan dengan profil pasien tertentu untuk menghasilkan penilaian awal yang dipersonalisasi mengenai jalan ke depan. Ini membantu dokter dan pasien merumuskan keputusan yang paling dekat dengan preferensi pasien berdasarkan hasil atau perilaku sebelumnya. Pandangan komprehensif terhadap penyakit-penyakit tersebut akan memungkinkan pasien untuk memiliki pemahaman yang lebih baik tentang lintasan kesehatan jangka panjang mereka dan untuk melakukan tindakan perbaikan antisipatif atau mengusulkan strategi pengobatan untuk mengatasi kondisi kesehatan mereka, termasuk pertimbangan kualitas hidup dan perawatan. Keragaman informasi ini memungkinkan pasien dan dokter untuk membuat keputusan yang lebih tepat ketika merencanakan perawatan dan pengobatan mereka di masa mendatang. Dokter dapat mengidentifikasi nuansa prognosis medis yang dihasilkan oleh AI seputar penyakit, prognosis, dan proyeksi harapan hidup, yang kemudian dapat dipetakan terhadap kondisi kesehatan individu, latar belakang, nilai-nilai, dan pertimbangan kualitas hidup. Dalam menawarkan saran yang spesifik dan kontekstual kepada pasien, dokter perlu membuat informasi ini mudah dipahami oleh pasien agar mereka dapat merumuskan keputusan. Hal ini karena informasi tersebut dapat berisiko menjadi berlebihan jika tidak disajikan dengan cermat.

Setelah pemahaman dan kapasitas mental pasien terbentuk, langkah selanjutnya adalah mencatat preferensi mengenai pengobatan dan perawatan medis di masa mendatang. Pemrosesan bahasa alami bermanfaat untuk mencari dan menemukan rekam medis dan untuk pengenalan suara guna merekam keputusan atau untuk mencatat kapan konsultasi dilakukan. Namun, catatan-catatan ini perlu diperiksa ulang dengan dokter untuk memverifikasi keakuratannya saat membuat keputusan akhir.

Aspek kedua dari pengambilan keputusan adalah tahap implementasi, di mana tiga pertimbangan penting muncul: apakah ada perubahan dalam hal pemikiran orang tersebut tentang keputusan dan keadaan pribadi serta ruang lingkup perawatan. Meskipun ruang lingkup perawatan dan keadaan pribadi dapat ditentukan, perubahan pikiran mungkin sedikit rumit. Ini karena ketika keputusan tersebut hendak diterapkan, pasien biasanya sudah tidak sadar, sehingga tidak dapat mengonfirmasi ulang keputusan tersebut. Hal ini khususnya menantang ketika keluarga menentang keputusan tersebut atau ketika keputusan tersebut bertentangan dengan penilaian klinis dokter mengenai perawatan. Penerapan AI mungkin tidak cukup bernuansa untuk menilai apakah pasien tertentu akan berubah pikiran karena sangat bergantung pada proses berpikir individu pada saat itu. Aspek ini akan lebih baik dikonfirmasi oleh keluarga pasien daripada AI.

Antusiasme terhadap aplikasi AI dan potensinya dalam pemberian layanan kesehatan mungkin telah menutupi kekhawatiran yang ditimbulkan oleh penggunaan dan keterbatasannya. Peringatan telah dikemukakan mengenai risiko inherennya yang timbul dari

bias, ketidakakuratan, agensi, tanggung jawab, akuntabilitas, dan kejelasan. Beberapa orang telah menyatakan kecenderungan yang terkendali untuk adopsinya, berhati-hati terhadap implikasi yang timbul dari akses data dan tata kelola, sementara yang lain mempertanyakan kemungkinan AI menggantikan fungsi manusia, yang tampaknya memberi bobot pada klaim bahwa kemampuan AI terlalu meningkat. Mekanisme pembelajaran mendalam menghadirkan kekhawatiran sosio-teknis yang memengaruhi proses pengambilan keputusan, yang mengarah pada masalah umum seperti hasil yang bias atau salah, korelasi yang dipertanyakan, dan keraguan transparansi. Kekhawatiran topikal yang relevan untuk memajukan pengambilan keputusan layanan kesehatan di mana integrasi AI berpotensi menimbulkan risiko dan kewajiban dieksplorasi di bagian selanjutnya. Tantangan-tantangan tersebut diklasifikasikan ke dalam dua aspek luas, yang pertama terkait dengan isu risiko terkait keandalan dan kepercayaan, dan yang kedua terkait dengan otonomi dalam pengambilan keputusan dan implikasinya terhadap hubungan dokter-pasien.

5.3 TANTANGAN, RESIKO, KEANDALAN, DAN TERPERCAYA

Integrasi AI dalam layanan kesehatan menghadirkan risiko bagi pengambilan keputusan klinis. Contohnya meliputi kesalahan medis, bias resep, dan interpretasi non-kontekstual dari hasil klinis, serta potensi variabel yang dihasilkan dari ketergantungan yang berlebihan pada analisis AI. Memang tepat jika konteks keseluruhan dianggap penting untuk menilai luaran ketika AI digunakan. Interpretasi non-kontekstual dalam pengambilan keputusan klinis penting karena berisiko mengabaikan atribut-atribut penting dan relevan dari masing-masing pasien, hubungan dengan keluarga yang lebih luas, komunitas, norma budaya, dan perbedaan geografis. Konfluensi faktor-faktor ini sering kali memengaruhi proses dan hasil pengambilan keputusan, di mana preferensi dapat berubah seiring waktu, seperti dalam kasus pengambilan keputusan lanjutan, yang menghadirkan ketidakpastian dengan keandalan keputusan. Masalah keandalan juga dipengaruhi oleh karakteristik AI dalam hal bias, ketika kumpulan data direplikasi di seluruh kelompok pasien yang serupa, yang mengakibatkan pendekatan menyeluruh yang tidak diinginkan untuk menasihati pasien. Dengan demikian, intervensi manusia diperlukan untuk menafsirkan hasil yang dihasilkan oleh AI untuk memastikan bahwa bias tersebut dihilangkan, dan bahwa setiap perbedaan yang jelas dari norma dan konteks diperbaiki untuk mencerminkan keputusan yang lebih selaras dengan nilai, norma, dan hubungan pasien.

Meskipun AI menawarkan prediksi tentang apa yang kemungkinan akan dipilih oleh orang-orang yang menderita kondisi kesehatan serupa, keputusan tersebut mungkin tidak selalu mencerminkan apa yang diinginkan pasien tertentu, karena pengambilan keputusan yang dibantu AI kurang bernuansa. AI memiliki keuntungan dalam menyuntikkan alasan dan keterampilan analitis ke dalam proses pengambilan keputusan tetapi mungkin kurang mampu memprediksi kebutuhan dan keinginan subjektif individu. Karena sifat dari hasil pengambilan keputusan lanjutan yang memengaruhi kesehatan secara keseluruhan dan keputusan yang mungkin menimbulkan risiko kematian, penerapannya harus mencakup pertimbangan tentang apa yang dipertaruhkan, karena konsekuensinya akan berbeda di berbagai bidang

pengambilan keputusan klinis, dengan ramifikasi yang berbeda. Faktor pembeda antara AI dan kecerdasan manusia adalah jumlah kebijaksanaan yang terlibat dalam pengambilan keputusan perawatan kesehatan. Sementara AI terkenal karena presisi, luas, dan kedalamannya, ia tidak memiliki nuansa kontekstual, penilaian nilai, intuisi, dan kebijaksanaan klinis, yang penting untuk pengambilan keputusan perawatan kesehatan. Lebih lanjut, AI tidak dapat mereplikasi pengalaman yang terakumulasi selama bertahun-tahun yang sebanding dengan pelatihan dokter. Tidak mengherankan bahwa kesesuaiannya untuk pengambilan keputusan perawatan kesehatan telah dipertanyakan. Kegunaan AI lebih dikenal dalam lingkungan yang beroperasi dalam jawaban yang jelas dan pasti; Akibatnya, kemampuannya mungkin lebih terbatas pada penyelesaian masalah non-dikotomis, yang sering ditemukan di dunia nyata. Pengambilan keputusan di bidang kesehatan sering kali melibatkan keseimbangan yang kompleks dan pertimbangan sosial yang sensitif dan saling bersaing. Macrae secara akurat menyoroti bahwa integrasi AI ke dalam layanan kesehatan menciptakan risiko baru dan menambah kekhawatiran yang sudah ada, yang menyebabkan bias yang semakin besar akibat berbagai asumsi yang 'tidak sensitif terhadap proses perawatan lokal'.

Masalah keandalan berkaitan erat dengan kepercayaan. Kepercayaan terhadap sistem layanan kesehatan yang terintegrasi dengan AI merupakan kekhawatiran utama yang dihadapi oleh para penyedia layanan kesehatan. Kepercayaan ini khususnya penting dalam tatanan layanan kesehatan yang mendasari hubungan dokter-pasien dalam memfasilitasi perawatan medis, pemulihan, penyembuhan, dan promosi kesehatan, yang merupakan bagian integral dari kebijakan dan etika kesehatan. AI dianggap tidak dapat dipercaya karena kurangnya kemampuan emosional dan ketidakmampuannya untuk bertanggung jawab atas tindakan yang biasanya memerlukan pertimbangan normatif tentang kepercayaan. Kekhawatiran tersebut tercermin dalam sebuah studi yang menemukan bahwa

para profesional layanan kesehatan cenderung tidak menggunakan AI dalam pemberian layanan kesehatan jika mereka tidak mempercayai teknologi tersebut atau memahami bagaimana AI digunakan untuk meningkatkan hasil pasien atau pemberian layanan yang spesifik untuk tatanan layanan kesehatan.

Sebuah studi serupa melaporkan pentingnya kepercayaan sebagai dasar untuk mengadopsi AI dalam perawatan kesehatan. Menghindari bahaya pasien adalah salah satu tujuan utama perawatan kesehatan; dengan demikian, ketika kepercayaan kurang dalam pengambilan keputusan perawatan kesehatan yang terintegrasi AI, langkah-langkah sangat penting untuk melindungi pasien terhadap risiko bahaya. Contoh perlindungan pasien terhadap risiko bahaya yang timbul dari sistem berbasis AI adalah perawatan sepsis, di mana penulis studi menganjurkan jaminan standar yang dinamis daripada statis yang mencerminkan sifat AI, dengan demikian mendukung setiap keputusan klinis dengan kewajiban untuk menghindari bahaya bagi pasien, memastikan perawatan pasien yang baik, dan mengamankan akuntabilitas kepada pasien.

Jelas bahwa keandalan dan kepercayaan adalah masalah yang terjalin erat dalam sistem perawatan kesehatan yang terintegrasi AI. Kedua gagasan tersebut berkontribusi pada masalah keamanan saat menilai hasil yang diprediksi oleh pembelajaran mesin dan apakah AI

telah memenuhi tujuannya dalam memfasilitasi pengambilan keputusan klinis. Kekhawatiran tentang kualitas dan keamanan pengambilan keputusan tersebut, yang timbul dari kemungkinan hasil yang salah atau bias, harus dimediasi oleh konfirmasi atau verifikasi manusia untuk menghilangkan ketidakpekaan terhadap preferensi dan keadaan aktual dari kasus tertentu tersebut. Lapisan perlindungan ini memungkinkan pendekatan manusia yang lebih intuitif terhadap proses pengambilan keputusan. Seperti yang disorot di atas, masalah sosio-teknis memerlukan pertimbangan jaringan manusia, sosial, dan organisasi tempat AI diterapkan. Meskipun seseorang mungkin menyarankan bahwa pengambilan keputusan perawatan kesehatan sebanding dengan AI dalam hal sifatnya yang secara intrinsik tidak stabil karena perubahan preferensi, kemajuan medis, dan risiko yang melekat, risiko dalam sistem perawatan kesehatan yang terintegrasi dengan AI bertambah ketika beberapa aspek berada di luar kendali dokter. Akibatnya, liabilitas yang timbul dari integrasinya ke dalam aplikasi layanan kesehatan menjamin pengembangan perlindungan dan langkah-langkah pencegahan untuk menghindari bahaya bagi pasien. Pengujian yang kuat dan berkelanjutan serta publikasi terbuka laporan keselamatan dapat membantu menumbuhkan kepercayaan dalam penggunaannya, serta konsultasi publik untuk mengidentifikasi penerimaan risiko tersebut dan cara-cara untuk memitigasinya guna membentuk pendekatan regulasi.

Kecerdasan merupakan fitur penting lainnya yang memengaruhi keandalan, keamanan, dan kepercayaan terhadap AI. Kecerdasan sering dikaitkan dengan sifat 'kotak hitam' AI dalam kaitannya dengan bagaimana AI mencapai keputusan dan hasilnya. Kekhawatiran ini muncul dalam pengambilan keputusan peradilan yang diotomatisasi AI, dengan kekhawatiran akan potensi diskriminasi putusan hukuman yang menyimpang dari konsep keadilan dan kewajaran. Dalam menghilangkan kekhawatiran ini, intervensi manusia sangat penting untuk memperbaiki ketidakadilan dan menghilangkan pengambilan keputusan mekanistik dan terstandarisasi yang memperkuat risiko sistemik dalam administrasi peradilan. Pendekatan serupa penting dalam mengatasi kekhawatiran tersebut dalam layanan kesehatan. Dalam pengambilan keputusan layanan kesehatan yang umum, pasien dan dokter terlibat dalam percakapan untuk memahami sifat penyakit dan mengklarifikasi kekhawatiran terkait pilihan pengobatan. Peluang ini mungkin lebih kecil kemungkinannya terjadi jika keputusan dihasilkan oleh AI karena akan ada kesulitan dalam memahami algoritma yang digunakan untuk pengambilan keputusan. Hal ini sangat penting karena keputusan layanan kesehatan berdampak pada kehidupan masyarakat, dengan potensi risiko yang muncul di masa mendatang. Beberapa cara untuk menumbuhkan kepercayaan dalam layanan kesehatan dalam mengurangi masalah 'kotak hitam' dalam AI meliputi memastikan integritas informasi, kompetensi penyedia layanan kesehatan, dan minat pasien untuk mendorong kepastian dalam pengambilan keputusan yang dibantu AI. Pilihan lain untuk mengungkap kejelasan AI adalah melalui hak atas penjelasan tentang proses pengambilan keputusan dalam Peraturan Perlindungan Data Umum (GDPR). Meskipun hak ini sebagian besar difokuskan pada keputusan otomatis, hak ini dapat diterapkan untuk menjelaskan algoritma pengambilan keputusan layanan kesehatan yang dibantu AI, seperti risiko, data layanan kesehatan, dan informasi yang memengaruhi pasien.

Keandalan, kepercayaan, dan kejelasan AI telah menimbulkan beberapa kekhawatiran penting dalam proses pengambilan keputusan. Aspek ini hanyalah sebagian dari gambaran besarnya. Kekhawatiran utama kedua berasal dari dilema etika otonomi dalam pengambilan keputusan dan implikasinya terhadap hubungan dokter-pasien.

5.4 TANTANGAN OTONOMI DAN RELASI DOKTER-PASIEN

Kekhawatiran etika seperti otonomi individu dan kepentingan terbaik sering kali tertanam dalam pengambilan keputusan di bidang kesehatan. Otonomi dalam pengambilan keputusan di bidang kesehatan menyiratkan bahwa setiap orang berhak untuk membuat keputusan yang sesuai dengan pandangan dunia mereka, meskipun belum tentu demi kepentingan terbaik mereka. Pengambilan keputusan di bidang kesehatan yang terintegrasi dengan AI mengundang pertimbangan tentang sejauh mana keputusan tersebut mencerminkan secara akurat pilihan otonom seseorang. Hal ini bergantung pada sejauh mana AI digunakan dalam proses pengambilan keputusan. Jika aplikasi AI utamanya bersifat asistif, maka kemungkinan besar otonomi individu dan penyedia layanan kesehatan tidak akan terganggu. Ketika AI diterapkan semata-mata atau secara ekstensif untuk memprediksi kemungkinan pilihan yang akan diambil seseorang, pertimbangannya menjadi lebih kompleks. Meskipun AI berharga untuk menghasilkan representasi profil kemungkinan keputusan pasien berdasarkan kriteria tertentu, deviasi biasanya terjadi pada tingkat individu. Oleh karena itu, preferensi pasien akan berbeda dari sampel umum, sehingga menghasilkan hasil yang berbeda dalam situasi tertentu. Pasien yang mampu mengungkapkan pilihan mereka akan mampu mengartikulasikan preferensi mereka, tetapi mereka yang menderita demensia atau berada dalam situasi terkunci mungkin tidak dapat menjalankan pilihan-pilihan ini. Mengadopsi pendekatan analitik rasional murni dalam membuat keputusan untuk pasien ini berisiko mengabaikan pertimbangan penting lainnya, seperti keinginan yang diungkapkan sebelumnya, pandangan orang tersebut, dan individualitas dalam menanggapi skenario semacam ini. Dalam beberapa kasus, intervensi yang diterapkan AI dapat berdampak negatif pada otonomi individu, di mana pilihan pasien dibatasi karena proyeksi risiko atau hasilnya diperkirakan tidak sesuai dengan kepentingan terbaik pasien.

Meningkatnya integrasi AI dalam perawatan kesehatan meningkatkan kemungkinan fungsionalitasnya menggantikan kendali manusia dan keahlian klinis yang memengaruhi hubungan dokter-pasien yang sudah mapan. Kekhawatiran tersebut berasal dari karakteristik AI yang tampaknya sempurna dibandingkan dengan manusia yang rentan terhadap kelelahan, terutama setelah berjam-jam konsultasi pasien atau prosedur pembedahan. Dalton-Brown, dalam mengeksplorasi masalah ini, menolak gagasan AI yang sepenuhnya menggantikan dokter, atas dasar bahwa langkah seperti itu akan merendahkan martabat perawatan kesehatan, karena penyembuhan dan kepedulian tetap menjadi fokus utama praktik perawatan kesehatan. Dokter lebih perseptif dalam analisis mereka tentang pilihan terapi dalam hubungan dokter-pasien dibandingkan dengan AI, yang berkontribusi terhadap pengembangan kepercayaan pasien dan peningkatan kualitas proses pengambilan keputusan melalui ekspresi kepedulian dan empati. Kemanusiaan umum tersirat dalam hubungan dokter-

pasien, yang mendukung otonomi pasien. Hubungan tersebut mewujudkan kepastian dan kepercayaan, yang melampaui diagnosis dan prognosis, menimbulkan keraguan tentang apakah AI dapat meniru empati semacam itu. Tidak sulit untuk membandingkan keterbatasan formula dan teknis AI dengan sentuhan manusia yang ditimbulkan oleh bahasa, ekspresi, dan gestur manusia.

Dalam merumuskan keputusan perawatan kesehatan yang akan diimplementasikan di masa mendatang, kesadaran akan nilai-nilai pasien dan kemampuan untuk mempertimbangkan prioritas yang saling bertentangan saat mengambil keputusan perawatan membutuhkan ketajaman medis, alih-alih kecerdasan semata. Penilaian etis dan moral lazim dalam proses pengambilan keputusan, dipupuk oleh rasa saling percaya, isyarat sosial tersirat, dan kasih sayang. Elemen-elemen ini hampir tidak dikatakan dimiliki oleh AI. Selain itu, meskipun aplikasi AI dapat dilatih untuk meningkatkan akurasi atau sensitivitas, masih belum pasti apakah AI dapat dilatih untuk mencerminkan artikulasi manusia. Sejauh ini, aplikasi AI dalam perawatan kesehatan bersifat fasilitatif, dan membantu mengurangi tekanan tugas administratif dan proses diagnosis yang lebih lanjut. Antarmuka manusia tetap menjadi fokus utama hubungan dokter-pasien.

Analisis di atas telah mengungkapkan area perhatian spesifik yang melibatkan pengambilan keputusan perawatan kesehatan tingkat lanjut, yang menyoroti pentingnya konteks dan intervensi manusia dalam proses pengambilan keputusan. Berbagai pertimbangan memengaruhi luaran layanan kesehatan, melampaui diagnostik, seperti nilai-nilai pribadi dan hubungan keluarga. Diskresi klinis, ketajaman, dan pengalaman tenaga kesehatan dalam mempertimbangkan berbagai pertimbangan yang saling bertentangan saat memberikan nasihat kepada pasien merupakan elemen penting yang belum tercakup secara memadai oleh aplikasi yang terintegrasi dengan AI. Akibatnya, intervensi yang terintegrasi dengan AI kemungkinan besar tidak akan menggantikan hubungan dokter-pasien, tetapi intervensi tersebut berharga dalam mendukung proses pengambilan keputusan, dan dalam mendekati pengambilan keputusan manusia. Bagian selanjutnya membahas opsi regulasi untuk melindungi pengambilan keputusan layanan kesehatan di mana AI terintegrasi ke dalam sistem.

5.5 PENGATURAN KEPUTUSAN DALAM LAYANAN KESEHATAN BERBASIS AI

Tantangan-tantangan yang diidentifikasi di atas menimbulkan pertanyaan tentang cara yang tepat untuk melindungi proses pengambilan keputusan ini dan memperbaiki masalah keamanan dalam kerangka kerja regulasi potensial. Hal ini secara khusus mengharuskan para pembuat undang-undang untuk terlibat dengan perkembangan AI dalam integrasi layanan kesehatan yang memengaruhi ranah pribadi kehidupan individu. Jelas bahwa terdapat keseimbangan yang rumit untuk dicapai ketika mengejar dorongan inovatif untuk memanfaatkan aspek-aspek AI yang menguntungkan bagi dokter dan pasien, serta mengelola risiko yang timbul dari intervensi AI. Para pembuat undang-undang menyadari bahwa pendekatan regulasi yang berbeda akan menghasilkan hasil yang berbeda; oleh karena itu, menjadi relevan untuk memeriksa akomodasi yang ada dalam undang-undang. Salah satu

contohnya adalah dengan memanfaatkan kerangka hukum yang berlaku dalam sistem gugatan perdata untuk isu-isu spesifik yang timbul dari masalah privasi dan kerahasiaan, sementara metode lain adalah dengan membuat undang-undang baru untuk mengakomodasi risiko tertentu tanpa adanya langkah-langkah tata kelola yang memadai. Yang pertama mungkin berarti respons yang lebih cepat dari undang-undang yang ada, yang diambil dari undang-undang Uni Eropa tentang hak asasi manusia, sementara yang kedua memungkinkan respons yang lebih spesifik terhadap risiko yang teridentifikasi, tetapi mungkin mengandung berbagai ketidakpastian karena sifat AI yang terus berkembang. Undang-undang tersebut juga dapat dianggap sebagai upaya untuk "mengejar ketertinggalan" dengan kemajuan AI. Setiap pendekatan memiliki kelebihan dan kekurangannya masing-masing.

Guihot dan rekan-rekannya menawarkan analisis yang bermanfaat tentang opsi regulasi yang tersedia untuk menanggapi masalah yang muncul dari aplikasi AI di berbagai bidang. Mereka menyarankan bahwa pemerintah mampu memengaruhi pemangku kepentingan industri dengan menetapkan ekspektasi arah yang ingin mereka kejar, sehingga membentuk tujuan dan bentuk regulasi sekaligus mendorong keterlibatan berkelanjutan dengan para pemangku kepentingan untuk mengidentifikasi potensi dan risiko. Para penulis juga menyukai pendekatan berbasis risiko multi-level terhadap regulasi. Pendekatan ini memiliki keuntungan memungkinkan pertumbuhan AI, sekaligus memungkinkan pemerintah untuk mempertahankan pengawasan atas kemajuannya. Akibatnya, regulasi apa pun yang dikembangkan akan lebih fleksibel untuk memediasi risiko yang ditimbulkan oleh AI.

Sebagian besar negara telah mengadopsi pendekatan 'tunggu dan lihat' yang hati-hati terhadap isu-isu AI yang muncul, sementara Uni Eropa telah menunjukkan inisiatif yang lebih proaktif dan berkembang dalam regulasi AI. Contoh relevan yang diambil dari Inggris, dan di tingkat Uni Eropa, Swedia berguna untuk menyoroti berbagai pendekatan untuk menanggapi penerapan AI di sektor kesehatan. Komisi Eropa telah mengembangkan berbagai inisiatif dalam konsultasi dengan para ahli industri dan di antara Negara Anggota untuk mengidentifikasi pendekatan regulasi potensial terhadap AI dan robotika, khususnya yang berkaitan dengan pengambilan keputusan otomatis dan mengakomodasi kepatuhan GDPR. Komisi mengadopsi pendekatan berbasis risiko yang permisif ketika berupaya memanfaatkan potensi AI sambil mengelola risiko yang timbul dari penerapannya, yang didukung oleh hak dan nilai UE, dan dipandu oleh AI yang bertanggung jawab dan etis. Pendekatan kolektif UE berpuncak pada Deklarasi Kerja Sama tentang Kecerdasan Buatan, dengan mayoritas Negara Anggota UE menjadi penandatangan Deklarasi tersebut pada tahun 2018. Buku Putih UE mengidentifikasi pendekatan berbasis risiko terhadap regulasi, memastikan risiko tinggi-rendah, dan area risiko potensial saat memetakan strategi khusus untuk menangani risiko ini tanpa menciptakan beban yang tidak semestinya pada para pemangku kepentingan. Pendekatan ini menanggapi berbagai tantangan risiko dan kebutuhan untuk menanamkan kepercayaan dalam intervensi AI, yang mendukung sikap yang berpusat pada manusia saat mengelola area sistem integrasi AI. Pendekatan semacam itu khususnya penting dalam sektor perawatan kesehatan, di mana terdapat risiko nyata yang melibatkan perawatan dan pengelolaan kesehatan masyarakat.

Inggris telah menerapkan pendekatan serupa dalam mengatur AI untuk tujuan transformasi ekonomi dan teknologi. Pusat Etika dan Inovasi Data yang dibentuk oleh pemerintah berfungsi sebagai penasihat bagi pemerintah, regulator, dan industri dalam pengembangan kerangka kerja AI. Kemitraan yang terjalin antara Google *DeepMind* dan rumah sakit NHS menandakan kolaborasi di sektor kesehatan untuk memberikan perawatan pasien yang terintegrasi dengan AI. Komite AI di House of Lords dengan tepat menunjukkan perlunya menyelesaikan aspek kejelasan ('kotak hitam') AI untuk menumbuhkan kepercayaan terhadap intervensi AI karena dampaknya yang besar terhadap kehidupan individu.

Swedia, sebagai Negara Anggota UE, merupakan penandatangan Deklarasi tersebut, dan dipandu oleh Strategi Nasional untuk AI dalam upaya memanfaatkan peluang yang dihadirkan oleh AI untuk memaksimalkan daya saingnya, mendorong kolaborasi publik-swasta, dan mengupayakan kerja sama aktif dengan negara-negara Nordik lainnya. Pendekatan Swedia ditujukan untuk mengamankan inovasi dan pengembangan sekaligus memitigasi risiko yang ditimbulkan oleh AI, yang didukung oleh kerangka kerja etika dan inklusivitas. Aplikasi AI telah diterapkan di berbagai sektor di Swedia, seperti layanan kesehatan, pendidikan, dan transportasi. Sektor layanan kesehatannya telah menyaksikan kemitraan lintas disiplin dan lintas sektor yang aktif, termasuk, antara lain, program Analytic Imaging Diagnostic Arena yang mempromosikan penelitian diagnostik pencitraan. Sikap ini menunjukkan strategi berbasis risiko tanpa menghambat pertumbuhan AI, dengan potensi manfaat bagi pembangunan yang etis dan berkelanjutan.

Contoh-contoh spesifik negara ini menggambarkan keseimbangan yang dicari ketika menangani risiko dan bias tertentu yang ditemukan dalam aplikasi terintegrasi AI. Akibatnya, mengklasifikasikan tingkat risiko memungkinkan risiko aktual diidentifikasi, memfasilitasi pemahaman yang lebih baik tentang bagaimana keputusan dicapai, dan bergerak selangkah lebih maju dalam memuaskan pertanyaan tentang transparansi dan akuntabilitas. Gagasan serupa tentang tata kelola digaungkan dalam analitik AI dalam analisis gambar otomatis dalam perawatan kesehatan. Ho dan rekan menyoroti bahwa 'sifat tata kelola biomedis semakin berbasis risiko, spesifik konteks, peka kasus, terdesentralisasi, kolaboratif dan dalam hal konstituen pengetahuannya, pluralistik'. Mengidentifikasi risiko dalam kerangka kontekstual dalam pengambilan keputusan perawatan kesehatan adalah langkah pertama menuju pemahaman isu-isu penting untuk menentukan tujuan dan penggunaan AI dalam proses klinis. Dalam mempertimbangkan jenis pendekatan regulasi yang akan diadopsi, penting untuk menghargai tujuan yang ingin dicapai. Seperti yang ditunjukkan di atas, keterlibatan industri adalah pendekatan yang disukai, karena sebagian besar penelitian dan pengembangan AI dilakukan oleh entitas swasta. Dengan demikian, menjadi penting untuk terlibat secara aktif dengan sektor ini untuk memaksimalkan potensi, sambil menjaga kesejahteraan, keselamatan, dan keamanan publik. Pemahaman yang lebih baik tentang risiko memungkinkan pengadilan dan pembuat undang-undang untuk mengidentifikasi cara terbaik untuk menanggapi risiko nyata ini dan tingkat tanggung jawab yang akan dikenakan sebagai bagian dari perlindungan regulasi.

AI secara bertahap terintegrasi dalam pemberian perawatan pasien, di mana pengambilan keputusan di bidang kesehatan tetap menjadi salah satu bidang utama di sektor kesehatan. Dalam konteks pengambilan keputusan lanjutan untuk perawatan kesehatan di masa mendatang, penerapan pendekatan berbasis risiko yang spesifik konsisten dengan kerangka regulasi yang luas yang dibahas di atas. Pendekatan semacam itu memfasilitasi pengendalian risiko yang lebih fleksibel dan akurat yang timbul dari sistem perawatan kesehatan yang terintegrasi dengan AI. Tujuan pengambilan keputusan perawatan kesehatan yang didukung AI adalah untuk membantu dokter dan pasien membuat keputusan yang lebih baik melalui peningkatan akurasi dan bukti terbaik yang tersedia. Meskipun bukan sepenuhnya mustahil bagi robot AI untuk melakukan konsultasi klinis dan mencatat keputusan perawatan kesehatan di masa mendatang, intervensi manusia sangat penting untuk mengendalikan dan memodulasi risiko apa pun yang timbul dari pendekatan yang dimodifikasi ini terhadap perawatan pasien. Misalnya, risiko yang timbul dari kesalahan transkripsi dalam konsultasi dokter-pasien melalui pemrosesan bahasa alami dapat dikurangi dengan prosedur verifikasi manusia lebih lanjut. Proyeksi lintasan penyakit dan harapan hidup dari analisis AI untuk pasien tertentu perlu dimediasi dengan pertimbangan kualitas hidup individu dalam menentukan perawatan kesehatan di masa mendatang. Aspek ini akan membutuhkan pengisian informasi yang cermat ke dalam sistem AI untuk memungkinkan hasil yang lebih terfokus dan berkualitas. Menghargai keterbatasan AI memungkinkan pengawasan yang lebih baik atas penggunaannya dalam pengambilan keputusan perawatan kesehatan. Tata kelola jangka panjang mencakup pengujian rutin aplikasi yang terintegrasi AI untuk menjaga keandalannya, dan akibatnya meminimalkan potensi kewajiban yang timbul dari penggunaannya.

Untuk membuat AI dapat dipahami oleh pasien dan dokter, sangat penting bagi mereka untuk dapat menelusuri kembali langkah-langkah pengambilan keputusan, bereksperimen dengan variasi yang ada seperti yang disajikan oleh analisis yang dibantu AI, dan memasukkan konteks penting ke dalam proses pengambilan keputusan. Salah satu pilihan untuk mencapai keputusan yang bertanggung jawab dan akuntabel adalah persetujuan berdasarkan informasi. Konsep persetujuan berdasarkan informasi merupakan ciri khas pengambilan keputusan perawatan kesehatan. Persetujuan juga didukung oleh kerangka kerahasiaan dan privasi yang lebih luas dalam GDPR dan perlindungan hak asasi manusia. Pasien harus dikonsultasikan dan persetujuan mereka diperoleh sebelum aplikasi yang terintegrasi AI digunakan dalam proses pengambilan keputusan. Mereka harus diinformasikan tentang tujuan penggunaan intervensi AI dalam prosesnya, risiko, efek, cakupan, jangkauan, dan keterbatasannya sehingga mereka terinformasikan saat membuat keputusan akhir terkait layanan kesehatan mereka di masa mendatang. Menyadari aspek ini dapat meredakan kekhawatiran terkait otonomi pasien dan memastikan bahwa nilai-nilai yang dirancang untuk dan dalam sistem AI benar-benar mencerminkan keinginan masyarakat. Pilihan-pilihan ini akan tunduk pada tinjauan yang disepakati antara dokter dan pasien sehingga mereka yakin akan keaslian keputusan yang dibuat yang didukung oleh AI. Proses ini memungkinkan dokter dan pasien untuk memastikan

keandalan keputusan yang diambil, dan jika keputusan tersebut dianggap merugikan, elemen-elemen tersebut dapat dihilangkan.

Ketidakpastian AI dan pengambilan keputusan lanjutan saling memengaruhi. Ketidakpastian yang ada dalam pengambilan keputusan lanjutan dimediasi oleh saran dan informasi yang diberikan melalui konsultasi dokter. Kejelasan AI yang lebih baik juga dapat diatasi dengan pemahaman algoritma AI. Dalam hal ini, dokter akan dihadapkan pada tantangan untuk mempelajari bagaimana AI memengaruhi pemberian saran kepada pasien dengan menjelaskan proses pengambilan keputusan yang dibantu AI. Risiko yang timbul dari tahap ini penting untuk ditangani, karena berpotensi memengaruhi kewajiban standar perawatan yang diharapkan. Fitur ini mengharuskan dokter untuk awalnya memahami mekanismenya dan kemudian menerapkannya untuk kepentingan pasien sehingga mereka dapat memahami informasi dan akhirnya membuat pilihan. Dalam konteks pengambilan keputusan lanjutan, ini berarti melindungi pilihan terlepas dari risiko yang ditimbulkan AI, tetapi mengambil langkah-langkah untuk mengurangi risiko tersebut. Para pembuat undang-undang perlu mempertimbangkan aspek ini ketika mengakomodasi risiko tambahan yang terlibat dalam pembuatan atau modifikasi regulasi.

Sifat AI yang terus berkembang telah menghadirkan kesulitan dan peluang regulasi. Pendekatan regulasi apa pun yang dipilih oleh negara tertentu, baik dengan menetapkan undang-undang khusus untuk AI atau memasukkan isu terkait AI di bawah cabang hukum yang ada, tetap menjadi tantangan tata kelola yang nyata. Sementara tantangan ini tetap ada, sama berharganya untuk mempertimbangkan menetapkan kode etik dan standar profesional untuk membantu penyedia layanan kesehatan dalam menangani isu-isu langsung yang muncul dalam proses pengambilan keputusan layanan kesehatan daripada secara otomatis menggunakan rancangan undang-undang. Contoh proposal untuk menetapkan kode etik dan komite etik ada di bidang regulasi robot perawatan (robot AI yang menyediakan perawatan untuk orang tua dan orang yang dibantu hidup). Kode etik dan harapan standar profesional menerangi kerangka tata kelola yang mencakup penyedia layanan kesehatan, produsen, dan lembaga pemerintah untuk menjaga kepercayaan dalam sistem layanan kesehatan dan mengurangi risiko terhadap otonomi, privasi, dan kerahasiaan pasien, yang pada akhirnya melindungi pasien dari bahaya. Salah satu hasil yang bermanfaat dari melembagakan kode dan standar ini adalah kemampuan untuk menawarkan jaminan kepada pasien dan mempromosikan kepercayaan bahwa keputusan perawatan kesehatan mereka dibuat dengan cara yang bertanggung jawab, dan bahwa dokter dipandu oleh norma-norma yang ditetapkan saat menjalankan kewajiban profesional mereka. Kedua nilai tersebut penting dalam menjaga hubungan terapeutik dokter-pasien.

Demikian pula, pelatihan berkelanjutan bagi penyedia layanan kesehatan dan edukasi pasien tentang aplikasi yang memengaruhi keputusan layanan kesehatan mereka sangat penting jika intervensi terintegrasi AI menjadi lebih umum dalam sistem layanan kesehatan. Menerapkan kebijakan yang tepat dan memastikan pelatihan yang memadai untuk penggunaan AI yang aman dan efektif merupakan cara untuk meminimalkan risiko yang timbul dari aplikasi AI dalam layanan kesehatan. Contoh pengurangan kelalaian medis adalah dengan

mengidentifikasi standar perawatan yang tepat dan mengambil langkah-langkah untuk mencegah terjadinya kesalahan. Strategi yang dilakukan meliputi pemeriksaan kualitas aplikasi AI yang terpasang, pemantauan dan pemeliharaan rutin aplikasi dan fitur keamanannya, mengidentifikasi risiko spesifik terhadap individu tertentu, dan selalu mendapatkan persetujuan dari pasien setiap kali AI akan diterapkan.

5.6 DAMPAK AI PADA SISTEM PELAYANAN KESEHATAN DI INDONESIA

Penerapan Artificial Intelligence (AI) dalam layanan kesehatan di Indonesia telah membawa transformasi signifikan, baik pada aspek teknis maupun manajerial. AI tidak hanya berfungsi sebagai alat bantu diagnostik, tetapi juga mendukung pengambilan keputusan klinis yang kompleks. Salah satu dampak utamanya adalah meningkatnya akurasi diagnosis melalui teknologi machine learning dan deep learning pada analisis citra medis. Contohnya, dalam radiologi, AI dapat mendeteksi kelainan paru-paru, kanker payudara, maupun stroke dengan lebih cepat dan tepat dibandingkan metode konvensional, sehingga mengurangi risiko kesalahan manusia dan mempercepat penentuan terapi. Selain itu, AI mendukung personalisasi layanan kesehatan dengan memberikan rekomendasi pengobatan sesuai kondisi unik setiap pasien. Sistem ini menganalisis rekam medis elektronik untuk mengidentifikasi riwayat penyakit, respons terapi sebelumnya, dan risiko komplikasi di masa depan, sehingga dokter dapat memberikan intervensi yang lebih presisi dan efisien. Pada level manajerial, AI digunakan untuk mengoptimalkan penjadwalan tenaga medis, memprediksi kebutuhan obat, serta mengelola antrian pasien, meningkatkan efisiensi operasional dan pengalaman pasien.

Meski demikian, penerapan AI masih menghadapi hambatan, seperti keterbatasan infrastruktur digital di daerah terpencil, kesenjangan literasi teknologi di kalangan tenaga medis, serta isu regulasi dan etika terkait perlindungan data pasien, akurasi algoritme, dan akuntabilitas kesalahan diagnosis atau rekomendasi terapi. Beberapa aplikasi nyata sudah mulai diterapkan di Indonesia, seperti platform telemedicine berbasis AI untuk konsultasi jarak jauh, penggunaan AI di radiologi dan laboratorium rumah sakit besar, serta integrasi dalam sistem asuransi kesehatan nasional untuk mendeteksi klaim ganda dan menganalisis pola penyakit populasi.

Dengan demikian, AI menghadirkan peluang besar sekaligus tantangan serius. Teknologi ini berpotensi meningkatkan kualitas layanan, efisiensi operasional, dan personalisasi perawatan pasien, tetapi keberhasilannya sangat bergantung pada kesiapan regulasi, infrastruktur digital, dan kompetensi sumber daya manusia. Sinergi antara pemerintah, institusi kesehatan, dan sektor swasta diperlukan agar penerapan AI memberikan manfaat maksimal tanpa mengabaikan keamanan, etika, dan hak privasi pasien.

5.7 KESIMPULAN

Bab ini telah menguraikan berbagai potensi sistem terintegrasi AI dalam pengambilan keputusan layanan kesehatan, dengan fokus pada pengambilan keputusan lanjutan. Sektor layanan kesehatan telah mengalami kemajuan, terutama dalam diagnostik dan prediksi penyakit, melalui teknologi jaringan saraf AI yang terus canggih dan pembelajaran mesin. Fitur

pemrosesan bahasa alami dan pembelajaran mendalam AI sangat berharga untuk memberikan informasi kepada tenaga kesehatan dan pasien saat membuat keputusan yang tepat. Pengambilan keputusan tingkat lanjut khususnya akan mendapat manfaat dari teknologi AI dalam memfasilitasi berbagai jenis perawatan dan rencana perawatan kesehatan di masa mendatang. Kami juga telah mempertimbangkan secara singkat bagaimana AI akan bermanfaat di masa pandemi untuk membantu penyedia layanan kesehatan yang berada di bawah tekanan yang sangat besar saat memberikan layanan kesehatan kepada masyarakat, terutama dalam membuat keputusan klinis yang sulit untuk melanjutkan atau menghentikan perawatan. Sifat AI yang terus berkembang, cair, dan tidak transparan membuatnya jauh dari sederhana, meskipun kemungkinannya tampaknya tak terbatas yang tersedia untuk dimanfaatkan demi kepentingan masyarakat. Kekhawatiran muncul mengenai risiko yang ditimbulkan oleh aplikasi tersebut, pertanyaan tentang otonomi pasien, konflik mengenai pengambilan keputusan, tingkat intervensi manusia, dan dampaknya terhadap hubungan dokter-pasien yang ada. Namun, diakui bahwa AI tidak dapat sepenuhnya menggantikan beberapa tugas dalam lingkup klinisi.

Penggunaan dan integrasi AI yang berkelanjutan dalam penyediaan layanan kesehatan akan terus memunculkan pertanyaan yang membutuhkan respons yang unik. Para legislator telah berupaya menanggapi isu-isu yang muncul, dengan serangkaian opsi regulasi yang memungkinkan untuk mengatur tantangan-tantangan ini dan memediasi risiko yang timbul dari penerapannya dalam layanan kesehatan. Pandangan yang lebih menguntungkan yang diadopsi adalah mengatur penggunaannya berdasarkan tingkat risiko yang ditimbulkan oleh AI, dan sesuai dengan konteksnya. Pertimbangan-pertimbangan ini mengungkapkan kekhawatiran penting yang dapat diatasi oleh legislasi untuk meminimalkan dampak AI pada kehidupan sehari-hari, terutama di sektor kesehatan yang semakin banyak digunakan. Pemahaman yang tepat tentang tantangan terkait AI serta isu hukum dan etika untuk pengambilan keputusan lanjutan dapat memandu pendekatan untuk lebih memahami penggunaannya, membantu dokter dalam menerapkan keputusan perawatan kesehatan masyarakat, dan bagi legislator untuk mengatur teknologi yang sedang berkembang ini dalam penyediaan layanan kesehatan.

BAB 6

KECERDASAN BUATAN DALAM PRAKTIK HUKUM

6.1 KEMAMPUAN AI DALAM PRAKTIK HUKUM SAAT INI

Hukum merupakan bagian integral dari cara masyarakat beroperasi. Profesi hukum sering dianggap tradisional dan hierarkis dalam arti bahwa struktur, etiket, dan keterampilan tertentu selalu memainkan peran penting. Namun, teknologi AI yang canggih telah memasuki profesi ini, memungkinkan analisis dan manajemen dokumen otomatis, analisis prediktif proses pengadilan, pemrosesan permohonan pengadilan, dan bahkan pengambilan keputusan. Hal ini menunjukkan bahwa AI telah diimplementasikan dalam berbagai aspek pekerjaan hukum, dan mengubah cara kerja pengacara. Salah satu area terbesar di mana AI telah diimplementasikan adalah analisis dan otomatisasi kontrak. Lebih dari itu, AI juga digunakan dalam analitik prediktif, administrasi peradilan, penelitian hukum, dan bidang lainnya. Tidak mungkin membahas secara rinci semua teknologi yang berbeda dalam satu bab dan, oleh karena itu, cakupannya di sini terbatas pada beberapa contoh utama, termasuk penelitian hukum, analitik kontrak, analitik prediktif, dan administrasi peradilan perdata dan pidana.

Riset Hukum

Pada dasarnya, riset hukum adalah tentang menemukan hukum dan sumber hukum yang relevan. Tanpa riset yang tepat, pengacara tidak akan mampu memberikan nasihat yang memadai kepada klien mereka. Khususnya dalam sistem hukum umum, menemukan sumber hukum yang tepat terkadang menantang karena doktrin preseden yudisial, yang menurutnya banyak prinsip hukum utama telah dibuat dan dikembangkan oleh hakim dari kasus ke kasus. Akibatnya, pengacara harus mampu menemukan preseden yang tepat dan, sayangnya, sistem pelaporan hukum di banyak tempat di seluruh dunia tidak selalu efektif. Di antara contoh perangkat lunak AI yang paling terkenal adalah yang disediakan oleh Westlaw dan LexisNexis, tetapi ada banyak perangkat lain di pasaran. Sebagian besar perangkat ini menyediakan analisis dokumen cerdas, di mana pengguna dapat mengunggah dokumen dan AI akan secara otomatis menyarankan otoritas yang relevan dengan dokumen tersebut tetapi tidak dikutip di dalamnya. Perangkat lunak ini juga akan mengungkapkan bagaimana pengadilan secara historis memperlakukan preseden yang relevan. Selain itu, mereka memberikan wawasan berbasis data, atau analisis litigasi, yang menyediakan informasi tentang hakim, pengadilan, ganti rugi, pengacara, firma hukum, dan jenis kasus untuk mengungkap pola tertentu dan membantu membangun strategi kasus. Oleh karena itu, alat berbasis AI ini berisi lebih dari sekadar mesin pencari untuk menemukan hukum atau kasus tertentu.

Analisis Kontrak

Hukum kontrak merupakan salah satu bidang hukum paling fundamental yang berlaku untuk setiap industri dan hubungan bisnis. Perjanjian yang memiliki kekuatan hukum tetap menimbulkan kewajiban dan perlindungan tertentu bagi para pihak yang terlibat.⁵ Penyusunan dan peninjauan kontrak bisa sangat memakan waktu, terutama jika kontrak

tersebut panjang, berisi puluhan atau bahkan ratusan halaman. Analisis dan peninjauan kontrak yang disempurnakan dengan AI dapat mempercepat proses secara signifikan, dan dapat membantu mengurangi dampak kesalahan manusia.

Ada puluhan perusahaan yang menawarkan analisis kontrak. Meskipun mereka menggunakan berbagai kode yang berbeda, dari sudut pandang pengguna, keseluruhan prosesnya hampir sama. Pertama, kontrak harus diunggah ke perangkat lunak. Kedua, pengguna harus menunjukkan apa yang perlu ditemukan dalam kontrak atau, dalam kasus platform yang lebih canggih, AI secara otomatis menghasilkan daftar aspek potensial yang akan ditinjau. Beberapa platform melangkah lebih jauh dengan menawarkan basis data manajemen kontrak, dan kemungkinan untuk bekerja lintas kontrak secara bersamaan. Sampai batas tertentu, program analitik kontrak masih membutuhkan intervensi manusia, dan peran utamanya adalah bekerja sama dengan penggunanya, untuk membuat keseluruhan proses lebih cepat, akurat, dan lebih hemat biaya.

Analisis Prediktif

Bukan hal yang aneh bagi para pihak yang bersengketa untuk bertanya kepada pengacara mereka tentang hasil realistis dari kasus mereka. Jawabannya seringkali bergantung pada berbagai faktor yang berbeda, termasuk bukti inti, metode hukum penyelesaian sengketa, yurisdiksi yang relevan, hakim, dan bahkan pengalaman pengacara yang terlibat. Namun, saat ini platform yang digerakkan oleh AI dapat memberikan intelijen dan analisis prediktif tentang proses dan hasil pengadilan sehingga pengacara dapat memperhitungkannya juga, di samping tebakan mereka yang sangat matang.

Setelah pengguna platform memasukkan semua fakta yang relevan dari kasus tersebut, AI membandingkannya dengan setiap kasus lain yang telah diajukan ke pengadilan. AI kemudian menghasilkan laporan dengan analitik prediktif yang mencakup berbagai aspek yang berbeda, termasuk: (a) perbandingan kasus dengan tren dan pola keseluruhan; (b) tingkat keberhasilan potensial jika kasus tersebut diajukan ke pengadilan; (c) kemungkinan hakim tertentu mengabulkan jenis mosi tertentu, (d) tipikal ganti rugi untuk jenis kasus tertentu, dan (e) kinerja masa lalu dan pengalaman litigasi pengacara. Ini merupakan informasi yang sangat berharga yang dapat membantu klien dan pengacara mereka membuat keputusan yang lebih tepat tentang cara menangani kasus tersebut.

Administrasi Peradilan Sipil

Contoh menarik dari administrasi peradilan modern disajikan oleh Pengadilan Resolusi Sipil (CRT) yang berpusat di British Columbia, Kanada. CRT adalah pengadilan daring yang dirancang untuk menyelesaikan empat jenis sengketa: (1) sengketa klaim kecil hingga CAD Rp. 50.000.000; (2) sengketa properti strata dengan jumlah berapa pun; (3) klaim kecelakaan kendaraan dan cedera hingga CAD Rp. 500.000.000; dan (4) sengketa perkumpulan dan koperasi dengan jumlah berapa pun.¹⁰ Meskipun sepenuhnya daring, CRT tidak sepenuhnya digerakkan oleh AI atau pengadilan yang sepenuhnya otomatis. Pengajuan ke CRT dilakukan melalui apa yang disebut "Solution Explorer", yaitu teknik inferensi (formulir tanya jawab) yang menggunakan bentuk dasar AI yang disebut sistem pakar. CRT menyajikan studi kasus penting untuk pengadilan dan pengadilan di masa mendatang yang lebih digerakkan oleh AI.

Contoh yang lebih maju adalah Courtal – sistem informasi pengadilan digital yang dirancang untuk Kementerian Kehakiman Estonia. Sistem ini memungkinkan administrasi elektronik berbasis AI untuk pekerjaan pengadilan Estonia di semua tingkatan. Sistem ini memungkinkan pendaftaran otomatis kasus, sidang, dan putusan, alokasi otomatis kasus kepada hakim, pembuatan panggilan, publikasi putusan, dan bahkan dukungan aktivitas juru sita. Selain itu, pada tahun 2021, Kementerian Kehakiman Estonia sedang menguji coba proyek “hakim AI” untuk mengadili sengketa klaim kecil dengan nilai kurang dari Rp. 70.000.000. Secara konsep, AI akan mengeluarkan keputusan yang dapat diajukan banding kepada hakim manusia. Diharapkan proyek ini dapat menyelesaikan penumpukan kasus dan membuat seluruh proses peradilan lebih tepat waktu dan hemat biaya.

Satu proyek terbaru yang layak disebutkan adalah SUPACE – Portal Mahkamah Agung untuk Bantuan Efisiensi Pengadilan, yang merupakan alat bantu berbasis AI yang diimplementasikan oleh Mahkamah Agung India. Sistem ini tidak dirancang untuk mengambil keputusan, melainkan untuk menganalisis data yang diterima dalam pengajuan perkara untuk mendukung pekerjaan para hakim. Jika berhasil di tingkat Mahkamah Agung, sistem ini dapat diimplementasikan di pengadilan yang lebih rendah untuk membantu menyelesaikan masalah penundaan dan manajemen perkara. Proyek ini resmi diluncurkan pada April 2021, sehingga pada saat penulisan ini, masih sedikit informasi yang tersedia tentang sistem ini. Namun, terungkap bahwa sistem ini mencakup chatbot yang dapat memberikan gambaran umum perkara, informasi tentang hukum yang relevan, dll. Sistem ini masih dalam tahap pelatihan oleh pengguna manusia.

ADMINISTRASI PERADILAN PIDANA

AI juga dapat digunakan dalam sistem peradilan pidana. Contoh paling terkenal dari penggunaan algoritma AI dalam perkara pidana meliputi: (1) Penilaian Keamanan Publik (PSA), (2) Instrumen Penilaian Risiko Pra-Persidangan (PTRA), (3) Profil Manajemen Pelanggar Pemasarakatan untuk Sanksi Alternatif (COMPAS), dan (4) Alat Penilaian Risiko Bahaya (HART).

PSA digunakan setelah penangkapan seseorang ketika hakim harus memutuskan apakah orang tersebut harus dibebaskan atau ditahan untuk menunggu persidangan. Sistem ini memprediksi kegagalan untuk hadir di pengadilan dan kemungkinan melakukan kejahatan lain. Dengan demikian, PSA menghasilkan skor yang dihitung berdasarkan beberapa faktor, termasuk usia, pelanggaran saat ini, dakwaan yang tertunda, pelanggaran ringan, kejahatan berat, atau hukuman kekerasan sebelumnya, kegagalan sebelumnya untuk hadir di pra-persidangan, dan hukuman penjara sebelumnya. PSA tidak menganalisis informasi pribadi yang sensitif, seperti ras, jenis kelamin, kebangsaan, agama, pendidikan, pekerjaan, alamat rumah, status perkawinan, dan lainnya. Sistem ini dikembangkan oleh Laura and John Arnold Foundation (LJAF) dan, digunakan di 19 negara bagian di Amerika Serikat. Mirip dengan PSA, PTRA adalah instrumen penilaian risiko yang dirancang untuk memprediksi risiko kegagalan untuk hadir di pengadilan dan penangkapan kriminal baru. Ini dikembangkan oleh Kantor Administratif Pengadilan AS (Kantor Layanan Praperadilan dan Masa Percobaan) dan digunakan oleh petugas layanan percobaan dan praperadilan Amerika Serikat. Perbedaan

utama antara PTRAs dan PSAs adalah algoritma penilaian yang mencakup berbagai faktor pengukuran. Secara khusus, PTRAs memperhitungkan usia terdakwa, riwayat kriminal, hukuman langsung, pencapaian pendidikan, status pekerjaan, kepemilikan tempat tinggal, masalah penyalahgunaan zat, dan bahkan status kewarganegaraan. Oleh karena itu, PTRAs mencakup beberapa informasi sensitif yang tidak dimiliki PSAs. Baik PSAs maupun PTRAs dipandang hanya sebagai alat pendukung, di samping faktor-faktor lain dan bukti yang tersedia bagi hakim, sehingga mereka dapat mempertimbangkan berbagai faktor dan membuat keputusan yang lebih tepat.

Seperti halnya semua teknologi semacam ini, baik PSAs maupun PTRAs memiliki pendukung dan penentang. Pendukung PSAs mengklaim bahwa algoritma ini dapat mengurangi jumlah orang yang ditahan di penjara sebelum persidangan. Akibatnya, mereka cenderung terhindar dari konsekuensi ekonomi dari penangkapan, termasuk kehilangan pekerjaan, hak asuh anak, dll. Lebih lanjut, hal ini seringkali menyebabkan berkurangnya penggunaan kondisi keuangan pembebasan (pembebasan dengan jaminan). Terlepas dari argumen-argumen ini, sistem tersebut tampaknya memiliki konsekuensi yang sangat besar. Khususnya, penentang menyuarakan kekhawatiran tentang cara prediksi dibuat. Hal ini karena prediksi didasarkan pada data dari banyak individu yang diklasifikasikan dalam risiko kelompok tertentu. Akibatnya, skor didasarkan pada sifat-sifat tertentu yang dimiliki oleh individu dalam kelompok yang berhasil atau gagal pada tingkat tertentu di masa lalu; oleh karena itu, skor tersebut tidak bersifat individual.

COMPAS dikembangkan oleh Northpointe Inc. dengan tujuan mengevaluasi risiko residivisme ketika hakim menentukan hukuman dalam kasus pidana. Algoritma ini menilai seseorang dari satu (risiko rendah) hingga sepuluh (risiko tinggi), dan skornya didasarkan pada jawaban atas 137 pertanyaan, baik yang dijawab langsung oleh terdakwa maupun yang diambil dari catatan kriminal, jika ada. Meskipun sistem ini hanya digunakan di beberapa yurisdiksi Amerika, sistem ini dikritik karena kurangnya transparansi dalam algoritmanya. Hal ini diakui oleh Mahkamah Agung Wisconsin dalam kasus *State v. Loomis*, di mana Mahkamah Agung mengakui bahwa skor risiko gagal menjelaskan bagaimana tepatnya data digunakan untuk menghasilkan hasil. Lebih lanjut, dinyatakan bahwa "kurangnya pemahaman pengadilan ini tentang COMPAS merupakan masalah yang signifikan dalam kasus ini".

HART adalah eksperimen Inggris yang dilakukan oleh Kepolisian Durham dan Universitas Cambridge. Sistem ini bertujuan untuk mengurangi pengulangan tindak pidana dengan mengidentifikasi dan membantu pelaku yang kemungkinan akan melakukan tindak pidana lagi dalam dua tahun ke depan. AI mempelajari keputusan yang dibuat antara tahun 2008 dan 2012 dan algoritmanya didasarkan pada sekitar 30 faktor, yang beberapa di antaranya tidak terkait dengan kejahatan yang dilakukan. Kritik utama terhadap sistem ini adalah fakta bahwa sistem tersebut mengategorikan individu ke dalam berbagai kelompok berdasarkan asal etnis, pendapatan, tingkat pendidikan, dan faktor-faktor lainnya.

6.2 REGULASI

Seperti di industri lain, terdapat perdebatan yang sedang berlangsung mengenai dampak AI terhadap beberapa fitur utama, seperti keamanan data, hak-hak asasi, dan banyak lagi. Ada beberapa instrumen yang menyediakan, baik secara langsung maupun tidak langsung, beberapa regulasi terkait AI. Sebagian besar dibahas dalam bab-bab lain volume ini, termasuk Peraturan Perlindungan Data Umum (GDPR) Uni Eropa, Buku Putih tentang AI, RUU terkait AI AS, atau Undang-Undang Jepang tentang Perlindungan Informasi Pribadi. Fokus utama mereka adalah mencapai berbagai kebijakan sosial, termasuk perlindungan data dan privasi, pencegahan bias atau diskriminasi, akuntabilitas, dan banyak lagi. Namun, semuanya membahas masalah AI secara umum, tanpa fokus khusus pada profesi hukum.

Salah satu instrumen yang didedikasikan untuk mengatur AI bagi pengacara adalah Piagam Etika Eropa tentang Penggunaan Kecerdasan Buatan dalam Sistem Peradilan dan lingkungannya, yang diadopsi pada Desember 2018 oleh Komisi Eropa untuk Efisiensi Keadilan (CEPEJ). Piagam tersebut mencakup lima prinsip inti, sebagai berikut: (1) Prinsip untuk menghormati hak-hak dasar; (2) Prinsip non-diskriminasi; (3) Prinsip kualitas dan keamanan; (4) Prinsip transparansi, ketidakberpihakan, dan keadilan; dan (5) Prinsip “di bawah kendali pengguna”. Prinsip pertama mengharuskan bahwa setiap penggunaan AI harus mematuhi hak-hak yang dijamin oleh Konvensi Eropa tentang Hak Asasi Manusia (ECHR) dan Konvensi tentang Perlindungan Data Pribadi. Inti dari prinsip ini adalah hak atas pengadilan yang adil. Prinsip kedua melarang penggunaan data yang dapat menyebabkan diskriminasi dalam bentuk apa pun terhadap individu atau kelompok. Selain itu, data tersebut tidak boleh mengarah pada analisis atau penggunaan deterministik. Prinsip ketiga mengharuskan perancang model pembelajaran mesin untuk menyertakan tim proyek multidisiplin, yang akan memanfaatkan keahlian praktisi hukum serta akademisi dari berbagai ilmu sosial. Setiap masukan ke sistem harus berasal dari sumber yang tersertifikasi juga. Prinsip keempat menyerukan keseimbangan antara kekayaan intelektual perusahaan yang menyediakan sistem berbasis AI dan kebutuhan akan transparansi, imparialitas, dan keadilan. Terakhir, prinsip kelima menyerukan otonomi pengguna, memberi mereka kemungkinan untuk meninjau data dan mengendalikan pilihan mereka.

Kelima prinsip ini sesuai dengan prinsip serupa yang mengalir dari standar perilaku profesional dan kode etik profesional. Peraturan ini memengaruhi cara pengacara dapat, atau seharusnya, menggunakan AI. Ada manfaat tertentu serta area yang perlu diperhatikan yang timbul dari meningkatnya penggunaan sistem AI di sektor hukum, khususnya oleh pengacara, advokat, dan hakim. Area ini mencakup, tetapi tidak terbatas pada, supremasi hukum, hak-hak fundamental, transparansi, dan keterjelasan.

Aturan Hukum

Terdapat beragam teori tentang aturan hukum, tetapi gagasan utamanya adalah tentang tata kelola pemerintahan yang baik, pembuatan hukum yang adil, dan penerapan hukum yang adil. Berbagai standar hukum dan kode etik profesi menyebutkannya sebagai salah satu prinsip utama, yang mewajibkan pengacara, advokat, hakim, dan profesional hukum lainnya untuk menegakkan aturan hukum. Beberapa teknologi, terutama yang digunakan

dalam administrasi peradilan, dapat memengaruhi elemen-elemen penting yang berfungsi dalam aturan hukum. Dengan mempertimbangkan hal tersebut, Organisasi untuk Kerja Sama dan Pembangunan Ekonomi (OECD) menetapkan dalam prinsip-prinsipnya tentang AI bahwa:

Para pelaku AI harus menghormati aturan hukum, hak asasi manusia, dan nilai-nilai demokrasi, di seluruh siklus hidup sistem AI. Ini mencakup kebebasan, martabat dan otonomi, privasi dan perlindungan data, non-diskriminasi dan kesetaraan, keberagaman, keadilan, keadilan sosial, dan hak-hak buruh yang diakui secara internasional.

OECD lebih lanjut merekomendasikan bahwa untuk tujuan ini, para pelaku AI harus menerapkan mekanisme dan perlindungan yang sesuai, termasuk kapasitas untuk penentuan nasib manusia. Pada bulan Juni 2019, G20 mengadopsi apa yang disebut Prinsip AI yang berpusat pada manusia dengan ketentuan yang sama persis yang diambil dari Prinsip-prinsip OECD.

Salah satu aspek kunci dari supremasi hukum adalah akses terhadap keadilan. Proyek-proyek yang digerakkan oleh AI seperti CRT Kanada, Courtal Estonia, dan berbagai teknologi berbasis cloud, membantu mewujudkan hak ini. CRT terbukti sangat berguna selama wabah COVID-19 ketika tetap buka dan beroperasi penuh 24 jam sehari, 7 hari seminggu. Pada saat yang sama, jutaan orang di seluruh dunia tidak memiliki akses yang layak terhadap keadilan, terutama di tempat-tempat di mana penguncian regional dan nasional diberlakukan. Akses terhadap keadilan merupakan masalah di banyak tempat bahkan sebelum pandemi, terutama karena kurangnya bantuan hukum. Meskipun selama pandemi banyak negara menawarkan sidang daring menggunakan sarana komunikasi audio-visual, penumpukan kasus di beberapa negara mencapai tingkat rekor. Contoh yang baik adalah Britania Raya, di mana, menurut House of Lords, pada Desember 2020 penumpukan kasus telah mencapai tingkat krisis, dengan lebih dari 8.000 orang ditahan menunggu persidangan. Oleh karena itu, tampaknya tepat untuk menyarankan agar ada investasi dan pengembangan berkelanjutan teknologi berbasis AI yang akan mendukung pekerjaan para pekerja di sektor hukum.

Hak-Hak Asasi

Alat analitik prediktif menjadi semakin populer karena memberikan intelijen yang sangat berguna, termasuk informasi tentang hukum yang relevan, catatan historis tentang ganti rugi yang diberikan, tetapi juga informasi tentang individu yang terlibat. Hal ini mengarah pada kemungkinan pembuatan profil hakim, pengacara, yurisdiksi, dan hal-hal lainnya. Di satu sisi, hal ini dapat merusak fungsi peradilan yang semestinya, tetapi di sisi lain hal ini membuat kegiatan publik lebih transparan dengan memungkinkan warga negara untuk mengetahui dan mengevaluasi hakim, pengacara, dan sebagainya. Praktik ini sudah dikenal di Amerika Serikat, tetapi telah dilarang di Prancis terkait dengan hakim dan juru tulis pengadilan. Kedua contoh yang saling bertentangan ini menunjukkan bahwa sikap terhadap teknologi ini sangat bervariasi.

Ciri umum lain yang sering dikaitkan dengan AI adalah potensi bias terhadap seseorang atau sekelompok orang. Berdasarkan contoh teknologi yang disebutkan sebelumnya, bias tersebut bisa berupa bias terhadap pengacara atau hakim, tetapi juga bias terhadap kelompok

tertentu, misalnya berdasarkan usia, jenis kelamin, atau karakteristik lainnya. Bias tersebut berasal dari input data ke dalam sistem, yang darinya AI belajar. Menghilangkan bias dari pembelajaran mesin merupakan area penelitian yang sedang berlangsung, tetapi pada prinsipnya, dimungkinkan untuk menghilangkan dan memodifikasi ciri input yang dapat menyebabkan bias. Masalah bias terkait erat dengan penghormatan terhadap hak-hak dasar. Kekhawatiran umum di berbagai industri adalah hak atas privasi dan larangan diskriminasi. Namun, dalam konteks profesi hukum, perhatian khusus memerlukan hak atas pengadilan yang adil, terutama jika kita memperhitungkan proyek "hakim AI" yang diujicobakan di Estonia atau sistem yang diterapkan di Mahkamah Agung India. Misalnya, chatbot yang memberikan ringkasan kasus tidak boleh menghasilkan hasil yang diskriminatif sehingga tidak ada pihak yang dirugikan, baik secara langsung maupun tidak langsung. Lebih lanjut, alat tersebut harus digunakan dengan menghormati independensi hakim dalam pengambilan keputusan. Dalam kasus AI yang memberikan keputusan aktual dalam suatu kasus, keputusan tersebut harus transparan dan dapat dijelaskan, serta harus tunduk pada pengawasan semua pihak yang terlibat, serta masyarakat umum.

Transparansi Dan Keterjelasan

Transparansi berkaitan dengan keakuratan dan relevansi informasi, yang memungkinkan seseorang untuk membuat keputusan yang terinformasi. Transparansi dalam konteks AI sebagian besar berkaitan dengan informasi tentang perusahaan dan platform AI yang ditawarkan, tetapi juga dapat berkaitan dengan klien yang ingin tahu bagaimana kasus mereka akan ditangani melalui penggunaan platform ini. Ketika AI melakukan tugas, manusia harus dapat memercayai hasil dan alasan di balik hasil tersebut. Transparansi juga menyiratkan kemungkinan untuk menjelaskan mengapa keputusan tertentu dicapai oleh AI, serta kemungkinan untuk menantangnya melalui penyelidikan resmi atau pengaduan yang akan memungkinkan revisi dan/atau peningkatan sistem/layanan. Yang terakhir ini sering disebut dengan istilah lain, yaitu "explainability" atau "AI yang dapat dijelaskan", yang melangkah lebih jauh dari transparansi. Masalah peraturan umum mengenai transparansi dan explainability telah dibahas dalam bab-bab sebelumnya dalam volume ini, khususnya dalam konteks GDPR dan Buku Putih UE. Namun, di luar ini, ada beberapa perkembangan yang dilakukan oleh Asosiasi Standar IEEE yang telah memulai serangkaian standar untuk perancang AI dan sistem otonom. Ini sedang dikembangkan dengan nama "seri IEEE P7000™", dan mencakup 13 standar terpisah untuk "masa depan sistem otonom dan cerdas yang selaras secara etis". Yang khususnya relevan dengan profesi hukum adalah Standar P7001 tentang Transparansi untuk Sistem Otonom. Tujuannya adalah untuk menyiapkan standar untuk mengembangkan teknologi otonom yang "dapat menilai tindakan mereka sendiri dan membantu pengguna memahami mengapa suatu teknologi membuat keputusan tertentu dalam situasi yang berbeda". Kelompok Kerja telah menetapkan lima kelompok yang akan mendapat manfaat dari standar tersebut, di antaranya adalah pengacara dan saksi ahli. Fokusnya adalah pada tingkat transparansi yang terukur dan dapat diuji, sehingga sistem otonom dapat dinilai secara objektif. Diharapkan bahwa proyek ini akan menawarkan cara untuk memberikan transparansi

dan akuntabilitas untuk sistem AI dalam profesi hukum dan profesi lainnya. Hingga Mei 2021, standar tersebut masih dalam pengembangan.

Kerahasiaan

Pengacara memiliki kewajiban untuk menjaga kerahasiaan urusan klien mereka. Selain itu, perlindungan khusus diberikan untuk semua komunikasi antara pengacara dan klien, serta dokumen yang disiapkan untuk litigasi, dan dokumen-dokumen tersebut tidak boleh diungkapkan kepada pihak ketiga. Yang terakhir ini umumnya dikenal sebagai "hak istimewa profesi hukum".

Pengacara terus-menerus berurusan dengan informasi sensitif, termasuk informasi identitas pribadi, data yang berkaitan dengan litigasi, komunikasi surel, dan banyak lagi. Oleh karena itu, mereka harus sangat waspada ketika mengadopsi teknologi baru. Banyak platform yang disempurnakan dengan AI berbasis cloud, artinya platform tersebut disimpan di jaringan server jarak jauh yang dihosting di internet. Hal ini mungkin berguna dalam situasi seperti wabah COVID-19 di mana mengakses infrastruktur TI di tempat dapat menjadi tantangan. Namun, merangkul AI dan konektivitas tinggi dapat membawa risiko ancaman siber dan pelanggaran data. Meskipun beberapa platform AI membanggakan diri dalam menawarkan keamanan tingkat bank dan enkripsi yang kuat untuk layanan di cloud dan di tempat mereka, bahkan itu tidak menjamin keamanan 100%. Oleh karena itu, penting untuk menyelidiki apakah penggunaan teknologi cloud memengaruhi kerahasiaan dan hak istimewa profesional hukum.

Di sebagian besar yurisdiksi, secara umum, pengacara dapat menggunakan teknologi berbasis cloud AI, meskipun ketentuan yang relevan bervariasi dalam detailnya. Misalnya, di Amerika Serikat, masalah ini ditangani oleh ketentuan yang ditetapkan oleh asosiasi pengacara negara bagian, dan sebagian besar dari mereka mengharuskan sistem cloud aman dan penyediaannya adalah organisasi yang memiliki reputasi baik. Namun, banyak dari mereka menyoroti bahwa kewajiban pengacara untuk melindungi tidak berakhir setelah penyedia yang memiliki reputasi baik dipilih; Langkah-langkah keamanan tersebut mencakup perawatan yang wajar dan berkelanjutan, dan tinjauan berkala. Di Eropa, Council of Bars and Law Societies of Europe (CCBE) membuat seperangkat pedoman yang mengharuskan pengacara untuk: (a) mempertimbangkan undang-undang perlindungan data dan prinsip-prinsip kerahasiaan profesional; (b) melakukan investigasi awal atas layanan, khususnya dalam hal pengalaman, reputasi, evaluasi teknis, akses luring, keamanan, dll.; (c) mengevaluasi infrastruktur TI internal dibandingkan dengan daring; (d) mempertimbangkan tindakan pencegahan dan transparansi kontraktual; dan (e) melakukan penilaian dampak. Pedoman ini mendahului penerapan GDPR, sehingga sekarang harus dibaca bersama dengan ketentuan GDPR. Secara khusus, perhatian khusus harus diberikan saat menggunakan platform AI yang melibatkan penyimpanan dan pemrosesan data pada server di negara lain. Data pribadi hanya dapat ditransfer ke negara ketiga dengan mematuhi ketentuan transfer data lintas batas yang ditetapkan dalam Bab V GDPR. Satu contoh lagi peraturan serupa datang dari Inggris dan Wales, di mana masalah kerahasiaan dan hak istimewa diatur oleh Standar dan Peraturan Solicitors Regulation Authority (SRA) dan Buku Pegangan Bar Standards Board (BSB). Mereka memerlukan

identifikasi, pemantauan, dan pengelolaan semua risiko. Selain itu, setiap informasi relevan yang dialihdayakan ke, dan dipegang oleh, pihak ketiga harus tersedia untuk SRA/BSB atau agen mereka untuk diperiksa.

6.3 KESIMPULAN

Dampak teknologi terhadap profesi hukum disoroti oleh American Bar Association dalam komentarnya terhadap Model Rules of Professional Conduct sebagai berikut: “seorang pengacara harus selalu mengikuti perubahan hukum dan praktiknya, termasuk manfaat dan risiko yang terkait dengan teknologi yang relevan (...)”. Berbagai contoh penerapan AI dalam layanan hukum menunjukkan bahwa AI dapat membantu pengacara meningkatkan keahlian dan pekerjaan mereka. AI dapat sangat membantu di masa seperti wabah COVID-19, dan umumnya ketika pekerjaan tertentu dapat atau harus dilakukan dari jarak jauh. Selain itu, pandemi memaksa pengadilan untuk lebih mempertimbangkan teknologi informasi dan komunikasi. Dalam konteks ini, Canadian CRT menyajikan studi kasus penting bagi pengadilan dan tribunal di masa depan yang lebih berbasis AI. Namun, peran alat-alat tertentu, terutama yang dapat memengaruhi pengambilan keputusan praperadilan, masih banyak diperdebatkan. Meskipun proses investasi dan pengembangan teknologi berbasis AI yang berkelanjutan perlu didorong, kolaborasi yang lebih besar diperlukan antara insinyur, peneliti, akademisi, dan sektor hukum. Semoga, dengan kolaborasi yang lebih besar, transparansi dan kemudahan dijelaskan pun meningkat. Para pelaku profesi hukum perlu memiliki pemahaman yang lebih baik tentang perangkat yang tersedia, kekuatan dan kelemahannya, dan khususnya fitur "prediktif"-nya. Contoh-contoh yang disajikan dalam bab ini menunjukkan bahwa hal-hal ini sangat penting terkait perangkat yang digunakan untuk administrasi peradilan. Menariknya, transparansi dan kemudahan dijelaskan juga menimbulkan pertanyaan tentang seberapa "otonom" sistem AI sebenarnya, dan sejauh mana seharusnya sistem tersebut. Tampaknya perlu ada interaksi yang seimbang di mana AI diperlukan untuk menghasilkan laporan kepada operator manusia beserta penjelasan tentang operasi dan keputusan AI.

Tingkat pengembangan perangkat AI sangat beragam, dan kekuatan serta kelemahannya telah diungkapkan. Mengadopsi AI dan konektivitas tinggi mungkin tampak tak terelakkan, tetapi juga dapat menimbulkan risiko lain, yaitu pelanggaran data. Pengacara selalu diharapkan untuk memenuhi standar tertinggi, dan prinsip ini juga berlaku untuk perlindungan data klien. Oleh karena itu, seiring dengan investasi dalam teknologi AI, investasi yang tepat dalam perlindungan keamanan siber harus mengikutinya. Hal ini tidak hanya mencakup sarana teknis untuk melindungi penyimpanan dan konektivitas data, tetapi juga harus melibatkan penerapan kebijakan yang relevan dan penyediaan asuransi siber. Meskipun semua langkah ini mungkin tidak mencegah pelanggaran/serangan siber (bisa dibilang tidak ada yang bisa), langkah-langkah ini telah menjadi standar industri. Karena pengacara sering berurusan dengan data pribadi dan rahasia yang sensitif, mereka harus tetap waspada dan berperan aktif dalam pemantauan berkelanjutan terhadap platform AI yang mereka miliki. Hal ini penting untuk perlindungan hak-hak dasar dan untuk memperkuat prinsip inti supremasi hukum.

BAB 7

AI DALAM PERLINDUNGAN HAK CIPTA MUSIK DIGITAL

7.1 EVOLUSI KE REVOLUSI MUSIK

Musik memainkan peran yang sangat penting dalam sejarah umat manusia. 'Ada empat tujuan musik yang jelas: tari, ritual, hiburan personal dan komunal, dan yang terpenting kohesi sosial, sekali lagi pada tingkat personal dan komunal.' Penemuan mesin cetak secara fundamental mengubah distribusi musik. Sejak musik cetak mesin pertama kali muncul pada abad ke-19, para komposer musik mulai mengandalkan reproduksi lembaran musik untuk menyebarkan musik mereka dengan cara yang jauh lebih cepat dan lebih nyata daripada sebelumnya. Kelahiran fonograf memungkinkan bentuk-bentuk baru rekaman musik dan distribusinya. Dengan penemuan piringan hitam, pertunjukan musik tidak lagi hanya bergantung pada musisi; kemudian, kaset menjadi format yang dominan; pada titik inilah pembajakan hak cipta memasuki panggung sejarah. Setelah kesuksesan CD pada tahun 1990-an, distribusi musik memulai perjalanan digital yang diikuti oleh peningkatan koneksi internet dan teknologi digital. Sementara musisi merayakan transmisi lintas batas yang mudah, mereka juga berjuang dengan cara mengendalikan (serta menerima remunerasi yang wajar untuk) musik mereka ketika pengunduhan dan streaming daring menjadi lebih populer. Bahkan, ada kekhawatiran yang meningkat tentang dampak hukum digitalisasi pada industri musik dan pihak-pihak terkait, karena dampak tersebut memengaruhi seluruh kelompok pencipta manusia.

Jika perubahan teknologi di atas diperlakukan sebagai revolusi dalam industri musik, AI akan memiliki dampak yang sama atau lebih besar. Kekhawatiran akan mesin yang menggantikan musisi manusia dalam penciptaan musik belum berhenti sejak Profesor David Cope di University of California di Santa Cruz mulai bereksperimen dengan komposisi algoritmik pada tahun 1981. Profesor Eduardo Miranda, salah satu peneliti terkemuka di bidang AI/musik, menunjukkan sikap yang lebih positif terhadap musik yang diproduksi AI. Dia menyebut AI sebagai 'sarana untuk memanfaatkan manusia daripada memusnahkannya ... pertimbangkan musik yang dihasilkan komputer sebagai benih, atau bahan mentah, untuk ... komposisi ... tidak begitu tertarik untuk memahami kreativitas dengan AI daripada ... tertarik pada AI untuk memanfaatkan ... kreativitas'. Namun, ini tidak mengurangi perdebatan tentang manusia vs mesin dalam pembuatan musik tetapi membuatnya lebih relevan, dan eksperimen terkait AI dalam membuat musik telah berkembang pesat sejak 2015. Misalnya, Melodrive memungkinkan perusahaan pengembangan game untuk membuat soundtrack khusus melalui sistem AI-nya, dan ini secara signifikan mengurangi biaya musik hingga 90%; Jukedeck adalah perusahaan rintisan yang berbasis di Inggris yang menggunakan AI untuk secara otomatis memproduksi musik. Pengguna hanya perlu memilih suasana hati, gaya, tempo, dan panjang maka sistem AI akan menghasilkan soundtrack yang cocok dengan elemen yang ditetapkan oleh pengguna (pengguna bisa mendapatkan lima lagu sebulan secara gratis). Banyak raksasa IT juga telah bereksperimen dengan proyek musik AI: perusahaan rintisan milik

Google, DeepMind memiliki proyek yang disebut WaveNet yang tujuannya adalah untuk menciptakan 'model generatif yang mendalam dari bentuk gelombang audio mentah ... mampu menghasilkan ucapan yang meniru suara manusia apa pun'; Proyek LNH Twitter mendorong pengguna untuk men-tweet bot LNH judul lagu dan bot kemudian men-tweet kembali trek instrumental pendek yang disusun oleh perangkat lunak; Proyek Google Magenta pada tahun 2016 bertujuan untuk menciptakan musik yang menarik dengan menggunakan sistem pembelajaran mesin; Di Tiongkok, raksasa TI Baidu mengembangkan komposer AI menggunakan perangkat lunak pengenalan gambar untuk mengubah gambar menjadi lagu; 哼趣 adalah perangkat lunak Tiongkok yang dapat menghasilkan musik berdasarkan pengguna yang menyenandungkan melodi pendek, dan pengguna juga dapat mengedit musik yang dihasilkan dengan memilih alat musik, gaya, dan durasi yang berbeda.

Semua hal di atas mengirimkan pesan yang jelas: sistem AI menjadi semakin populer untuk pembuatan musik dan mau tidak mau bersaing dengan musisi manusia sampai batas tertentu. Diprediksi bahwa di masa depan AI akan memainkan peran penting di sepanjang tiga lini dalam industri musik menurut perkembangan aplikasi musik AI: hiburan mandiri, simulasi lingkungan, dan pembuatan musik asisten dan independen.

Pertama-tama, aplikasi AI, termasuk Amper Score, Jukedeck, dan 哼趣, menjadikan pembuatan musik sebagai tugas satu menit yang mudah bagi pengguna publik yang menggunakan dan berbagi musik yang diproduksi AI ke media sosial untuk tujuan bersenang-senang. Intervensi pengguna publik dalam pembuatan musik diminimalkan karena musik yang dihasilkan sangat bergantung pada operasi mandiri dan desain algoritma aplikasi AI. Namun, sulit untuk mengatakan apakah AI harus dianggap sebagai pencipta musik yang dihasilkan karena kemampuan kreatif sistem AI tampaknya diragukan. Kedua, tanpa membayar biaya lisensi hak cipta yang lebih tinggi kepada musisi manusia, musik latar yang diproduksi AI berbiaya rendah atau musik yang disimulasikan AI di lingkungan sekitar dapat menempati sebagian besar industri hiburan, termasuk restoran, pub, dan tempat umum lainnya. Perusahaan pengembang gim atau proyek jangka panjang lainnya dengan anggaran terbatas juga dapat menggunakan musik AI untuk mengurangi biaya. Melodrive dapat 'menciptakan aliran musik orisinal yang tak terbatas dan bervariasi secara emosional secara real-time – idenya adalah bahwa musik tersebut beradaptasi dengan apa yang terjadi di dalam gim pada titik waktu tertentu'. Kasus ini memberikan contoh yang baik tentang apakah musik latar ini, yang diproduksi oleh sistem AI, harus dianggap sebagai karya yang dilindungi hak cipta, dan oleh karena itu apakah sistem AI harus diperlakukan sebagai penciptanya. Terakhir, AI juga dapat membantu musisi manusia dalam penciptaan musik, atau bahkan AI dapat dilibatkan sebagai kolaborator kunci dalam penciptaan musik. Startup Popgun yang berbasis di Australia mengumumkan bahwa produk AI mereka, Splash Pro, adalah alat untuk 'memberdayakan jutaan musisi untuk menemukan, terinspirasi, dan mengekspresikan kreativitas mereka'. Kolaborasi, alih-alih menggantikan masukan manusia, adalah tujuannya. Namun, hal ini akan menimbulkan pertanyaan tentang siapa pencipta karya tersebut. Karya tersebut dapat dianggap sebagai karya hibrida yang dibuat oleh manusia dan sistem AI. Lebih lanjut, ada kemungkinan juga bahwa sistem AI suatu hari nanti dapat menciptakan musik secara mandiri

dengan intervensi manusia yang minimal. Oleh karena itu, bagaimana menangani musik yang diciptakan secara mandiri oleh AI masih menjadi perdebatan.

Namun, sifat artifisial sistem AI dan apa yang disebut 'proses penciptaan' mereka berada di luar pandangan dan interpretasi tradisional tentang apa yang mendefinisikan karya yang dapat dilindungi hak cipta dan apa yang berkaitan dengan kreativitas dalam konteks hak cipta; hal ini memengaruhi dasar-dasar sistem hukum hak cipta yang ada dan dapat mengakibatkan evaluasi ulang total tentang apa yang seharusnya dilindungi oleh hukum hak cipta.

7.2 LEGALITAS MUSIK BERBASIS KECERDASAN BUATAN DI INDONESIA

Penggunaan kecerdasan buatan atau Artificial Intelligence (AI) dalam proses penciptaan musik digital membawa dampak yang cukup besar terhadap sistem perlindungan hak cipta di Indonesia. Teknologi ini tidak lagi sekadar berfungsi sebagai alat bantu, tetapi telah mampu menghasilkan karya musik secara mandiri tanpa adanya campur tangan manusia dalam tahap kreatif yang signifikan. Kondisi tersebut menimbulkan pertanyaan mendasar mengenai keabsahan karya yang dihasilkan AI, sebab dalam peraturan yang berlaku saat ini, khususnya Undang-Undang Nomor 28 Tahun 2014 tentang Hak Cipta, pencipta masih dipahami secara tegas sebagai seorang manusia atau individu. Artinya, karya musik yang sepenuhnya lahir dari sistem AI berpotensi tidak dapat memperoleh perlindungan hukum yang sah karena tidak memenuhi kriteria pencipta sebagaimana yang telah ditentukan.

Namun, permasalahan tidak berhenti sampai di situ. Ketika AI digunakan dalam bentuk kolaboratif dengan manusia, misalnya sebagai instrumen kreatif untuk menyusun aransemen, menciptakan melodi, atau menghasilkan efek suara tertentu, muncul persoalan baru tentang siapa yang berhak diakui sebagai pencipta utama dan bagaimana pembagian hak moral maupun hak ekonomi harus ditentukan. Di satu sisi, keterlibatan manusia tetap penting, tetapi di sisi lain kontribusi AI juga tidak bisa diabaikan. Hal ini memperlihatkan adanya ruang abu-abu yang perlu segera dijawab oleh hukum di Indonesia.

Contoh nyata dapat dilihat dari fenomena penggunaan AI dalam industri musik populer di Indonesia, misalnya ketika beberapa musisi independen menggunakan aplikasi berbasis AI untuk membuat melodi dan aransemen lagu yang kemudian diunggah ke platform digital seperti YouTube dan Spotify. Dalam beberapa kasus, karya tersebut berhasil menarik banyak pendengar, tetapi kemudian muncul perdebatan mengenai siapa pemegang hak cipta yang sah: apakah musisi yang hanya memberi instruksi dan menyempurnakan hasil AI, ataukah perusahaan pengembang teknologi AI yang menyediakan sistem pencipta lagu tersebut. Situasi ini bahkan pernah menimbulkan polemik di kalangan kreator musik lokal, karena ada kekhawatiran bahwa karya musik yang dihasilkan AI bisa menggeser peran pencipta manusia sekaligus mempersulit proses perlindungan hak cipta jika tidak diatur dengan jelas.

Perkembangan ini menunjukkan bahwa kehadiran AI menuntut adanya pembaruan regulasi hak cipta agar mampu memberikan perlindungan hukum yang jelas terhadap karya musik digital, baik yang sepenuhnya dihasilkan oleh AI maupun yang tercipta melalui kolaborasi antara manusia dan mesin. Tanpa adanya pembaruan, akan terus timbul

ketidakpastian hukum, terutama terkait dengan kepemilikan, distribusi, dan legalitas komersialisasi karya musik digital yang melibatkan teknologi kecerdasan buatan. Dengan kata lain, regulasi yang ada harus mampu mengikuti perkembangan teknologi agar hak-hak kreator tetap terlindungi dan kepastian hukum tetap terjaga.

7.3 KEMAMPUAN AI UNTUK KREATIVITAS

Tanpa definisi AI yang disepakati secara universal, sulit untuk memahami sifat sistem AI dan kemampuannya untuk kreativitas, serta menentukan apakah keluaran yang dihasilkannya memenuhi syarat untuk perlindungan hak cipta. Sebuah presentasi yang disampaikan oleh John Launchbury, Direktur Information Innovation Office (I2O) di Defense Advanced Research Projects Agency (DARPA) di Departemen Pertahanan Amerika Serikat, membagi penelitian AI menjadi tiga gelombang berdasarkan karakteristiknya, yaitu 'pengetahuan buatan', 'pembelajaran statistik', dan 'adaptasi kontekstual'. Oleh karena itu, mungkin dimungkinkan untuk mengeksplorasi kemampuan sistem AI melalui pemahaman penelitian AI pada berbagai tahap dan dengan demikian menangkap keluaran yang dihasilkan oleh sistem AI dan mengujinya untuk tanda-tanda kreativitas.

Didukung AI

Sistem AI yang didasarkan pada 'pengetahuan buatan' tidak memiliki kemampuan belajar, penanganan ketidakpastian yang buruk, dan hanya dapat menangani masalah yang didefinisikan secara sempit. Sistem AI semacam itu harus diperlakukan sebagai mesin yang mengikuti aturan yang ditentukan oleh manusia dan mendukung manusia untuk menyelesaikan suatu tugas. Oleh karena itu, musik yang dihasilkan AI tersebut seharusnya disebut musik yang didukung AI, merangkul intervensi penuh manusia dan di bawah kendali substansial manusia, karena sistem AI hanya menjalankan aturan yang ditetapkan manusia untuk menghasilkan musik. Hal ini serupa dengan manusia yang menetapkan aturan dalam program perangkat lunak di mana angka-angka dimasukkan ke dalam perangkat lunak, menghasilkan hasil berdasarkan aturan yang telah ditetapkan sebelumnya. Perbedaannya adalah sistem AI dapat menangani jumlah yang lebih besar dan aturan yang lebih rumit. Jelas, tidak ada kreativitas di sisi sistem AI dan, oleh karena itu, dari perspektif hak cipta, tidak ada keraguan tentang kepengarangan manusia atas musik yang dihasilkan oleh sistem AI.

Berbantuan AI

Dalam psikologi, kebaruan dan kesesuaian merupakan dua komponen penting untuk mengidentifikasi kreativitas – artinya, produk atau proses yang dihasilkan dari suatu aktivitas kreatif bersifat baru dan mengandung nilai. Namun demikian, definisi ini ambigu ketika menilai derajat 'baru' dan 'nilai'. Sebuah opini meragukan bahwa aktivitas mesin memungkinkan untuk mengkualifikasi elemen kreativitas apa pun, karena meyakini bahwa kreativitas 'sering dipandang secara spiritual, dipandang sebagai satu-satunya karakteristik 'manusia' yang unik, yang mendefinisikan suatu area pengalaman ... berpikir kreatif adalah benteng martabat manusia di zaman di mana mesin, terutama komputer, tampaknya mengambil alih aktivitas keterampilan rutin dan pemikiran sehari-hari'.

Faktanya, sistem AI tidak hanya mengambil alih aktivitas keterampilan rutin. Gelombang kedua penelitian AI adalah buktinya. Sistem AI semacam itu didasarkan pada 'pembelajaran statistik' dalam presentasi Launchbury. Berkat perkembangan pesat big data dan daya komputasi, serta sistem algoritma yang lebih baik (tiga pilar kesuksesan sistem AI), sistem AI saat ini dapat mendefinisikan aturan sendiri melalui pengelompokan dan pengklasifikasian kumpulan data besar, lalu menggunakan aturan yang telah ditentukan ini untuk memprediksi dan membuat keputusan sendiri sampai batas tertentu. Inilah mengapa manusia berpikir mereka dapat "berbicara" dengan Siri Apple dan bagaimana aplikasi AI dapat "menciptakan" musik. Dalam konteks sains manusia, istilah "kreativitas" tampak misterius karena sering dijelaskan dengan gagasan yang samar seperti "inspirasi" dan "intuisi", tetapi gagasan ini menunjukkan bahwa kreativitas membutuhkan sumber imajinasi. Manusia dapat membayangkan hal-hal yang belum pernah mereka lihat; namun, "imajinasi" sebagian besar komposer AI terbatas atau minimal, jika dibandingkan dengan kebebasan dan fleksibilitas imajinasi manusia yang tak terbatas.

Kenyataannya, sistem AI masih berbasis pembelajaran statistik, yang berarti sistem AI tidak dapat mempelajari dan memprediksi secara imajinatif berbagai hal dengan fitur di luar data yang disediakan manusia – AI hanya memiliki kemampuan pengambilan keputusan yang terbatas dan kemampuan prediksi yang terbatas, alih-alih kebebasan penuh atas kemampuan tersebut (seperti manusia). Selain itu, sistem AI tersebut tidak memahami dan menjelaskan data, melainkan menganalisis fitur-fitur data yang ada. Sebagaimana pernyataan MuseNet di situs webnya, "MuseNet tidak diprogram secara eksplisit dengan pemahaman musik, melainkan melalui pola harmoni, ritme, dan gaya yang ditemukan dengan belajar memprediksi token berikutnya dalam ratusan ribu berkas MIDI". Lebih lanjut, kemampuan pengambilan keputusan sistem AI yang terbatas bergantung pada pengaturan algoritma oleh insinyur manusia. Sementara insinyur manusia menetapkan parameter abstrak dalam sistem algoritma, komposer AI dapat menghasilkan komposisi musik yang berbeda ketika pengguna memasukkan kondisi yang sama; mereka akan menghasilkan hasil yang sama jika mereka menetapkan parameter konkret dalam algoritma. Intervensi manusia yang nyata muncul di setiap tahap proses produksi (mulai dari data pelatihan yang dipilih oleh manusia hingga pengaturan algoritma yang menentukan tingkat pengambilan keputusan AI). Dalam hal ini, kontribusi manusia menjadi pusat perhatian dan sistem AI hanya berstatus asisten. Oleh karena itu, sudah sepantasnya jika keluaran yang dihasilkan oleh sistem AI tersebut disebut keluaran berbantuan AI. Namun, mengingat aplikasi musik AI (sistem AI gelombang kedua) dapat menghasilkan musik yang tidak dapat diprediksi oleh manusia, bahkan dengan kemampuan pengambilan keputusan mereka yang terbatas, elemen-elemen yang tidak dapat diprediksi tersebut dapat dilihat sebagai bukti kontribusi material dari aplikasi AI. Oleh karena itu, patut dipertimbangkan apakah hal ini dapat menyesuaikan tingkat kreativitas tertentu, tetapi hal ini mungkin sangat bergantung pada bagaimana hukum hak cipta yang berlaku menafsirkan istilah 'kreativitas' dan apakah interpretasi tersebut dapat diperluas ke situasi pengambilan keputusan yang terbatas ini.

Dihasilkan Oleh AI

Dalam laporannya tentang Percakapan tentang Kekayaan Intelektual dan Kecerdasan Buatan, Organisasi Kekayaan Intelektual Dunia (WIPO) mendefinisikan 'dihasilkan oleh AI' sebagai 'penghasilan keluaran oleh AI tanpa campur tangan manusia'. Hal ini serupa dengan gelombang ketiga sistem AI Launchbury yang berbasis pada 'adaptasi kontekstual'. Sistem AI semacam itu belajar lebih banyak dari pemahaman fenomena dunia nyata daripada dari data, dengan kemampuan menjelaskan, menalar, dan mengabstraksi informasi yang dipersepsikan, untuk mengembangkan makna baru dengan sendirinya. 'Mengembangkan makna baru' menyiratkan penciptaan karena hal ini menunjukkan bahwa sistem AI sedang 'berpikir': mereka mencoba mengintegrasikan informasi yang dipersepsikan, kemudian sistem AI akan menjelaskan dan menalar informasi ini. Lebih penting lagi, mereka mengabstraksi dan mengeksplorasi hal-hal baru di luar informasi tersebut. Kemampuan mengabstraksi dan mengeksplorasi hal-hal baru dari informasi yang ada, dari sudut pandang penulis, pada akhirnya mengarah pada aktivitas kreatif. Tidak diragukan lagi bahwa kreativitas ada dalam keluaran yang dihasilkan. Menurut penulis, keluaran yang dihasilkan oleh sistem AI tersebut seharusnya disebut keluaran yang dihasilkan AI, bukan keluaran yang diciptakan AI, untuk membedakannya dari ciptaan manusia. Hal ini karena sifat sistem AI tersebut dan keluaran yang dihasilkannya berada di luar cakupan sistem hak cipta yang ada, yang dirancang khusus untuk ciptaan manusia.

7.4 TANTANGAN HAK CIPTA MUSIK BUATAN AI

Kepengarangan

Situs web resmi WIPO memberikan pengantar singkat tentang Konvensi Berne untuk Perlindungan Karya Sastra dan Seni. Dinyatakan bahwa Konvensi Berne 'berurusan dengan perlindungan karya dan hak-hak penciptanya'. Oleh karena itu, apakah istilah 'pencipta' mencakup makhluk non-manusia atau tidak akan menghasilkan hasil yang sama sekali berbeda: perlindungan atau non-perlindungan. Sayangnya, Konvensi Berne tidak mendefinisikan kepengarangan, mengingat perkembangan teknologi agak berbeda pada tahun 1886, yaitu saat Konvensi diadopsi, hal ini tidak mengherankan. Definisi atau interpretasi istilah 'pencipta' mungkin dianggap tidak perlu karena logika umum mengaitkan istilah ini dengan orang yang menciptakan karya tersebut. Pada awal tahun 1992, Profesor Ricketson membela pendapat tentang kepengarangan manusia dalam Konvensi Berne. Faktanya, terlepas dari aturan yang menyatakan bahwa Konvensi memberikan hak moral kepada pengarang (rasanya aneh untuk mengatakan mesin memiliki hak moral) atau ketentuan perlindungan yang mengatur masa hidup pengarang ditambah periode waktu setelah kematiannya, aturan-aturan tersebut menunjukkan bahwa Konvensi Berne mempertahankan gagasan tentang kepengarangan dan hak pengarang yang berpusat pada manusia. Dari sudut pandang ini, identitas kepengarangan AI sulit dijelaskan di tingkat internasional. Namun, isu ini tampak jauh lebih jelas dalam undang-undang hak cipta nasional.

Kantor Hak Cipta AS memberikan interpretasi yang jelas tentang kepengarangan yang memungkinkan pencipta untuk 'mendaftarkan karya asli kepengarangan, asalkan karya

tersebut dibuat oleh manusia'. Ini menyusul perselisihan pelanggaran potret diri monyet nera. Dalam kasus *Naruto v Slater*, seekor monyet nera mengambil serangkaian foto dirinya sendiri sementara David Slater, seorang fotografer profesional, meninggalkan kameranya di cagar alam. Fotografer itu menerbitkan sebuah buku termasuk potret diri monyet nera, tetapi kepengarangannya diperdebatkan oleh People for the Ethical Treatment of Animals, yang mengklaim bahwa monyet harus menjadi pemilik hak cipta dari potret diri tersebut. Dalam *Naruto*, isu utamanya adalah apakah seekor hewan dapat mengklaim hak cipta dan mengajukan klaim pelanggaran hak cipta berdasarkan hukum hak cipta AS. Akhirnya, Pengadilan Banding Kesembilan memutuskan bahwa monyet makaka tidak memiliki kedudukan konstitusional maupun hukum untuk mengklaim pelanggaran hak cipta. Hal ini karena hakim menemukan kata-kata yang digunakan dalam undang-undang hak cipta 'menyiratkan kemanusiaan dan tentu saja mengecualikan hewan yang tidak menikah dan tidak memiliki ahli waris yang berhak atas properti secara hukum'.

Undang-undang hak cipta Tiongkok juga mengklarifikasi bahwa 'karya warga negara Tiongkok, badan hukum, atau organisasi lain, baik yang diterbitkan maupun tidak, akan menikmati hak cipta sesuai dengan Undang-Undang ini'. Teknologi saat ini tidak dapat menciptakan sistem AI yang cukup cerdas untuk membuat kontrak atau mengklaim kemampuan hukum untuk menuntut atas pelanggaran haknya sendiri. Oleh karena itu, sistem AI tidak dapat diakui sebagai badan hukum atau organisasi lain. Dalam praktiknya, Pengadilan Internet Beijing memberikan interpretasi yang jelas tentang kepengarangan manusia pada bulan April 2019 dalam putusannya dalam kasus *Feilin v Baidu*. Salah satu perselisihan, dalam kasus ini, adalah apakah laporan tentang analisis yudisial industri film dan hiburan yang dihasilkan secara otomatis oleh perangkat lunak bernama Wolters Kluwer China Law & Reference merupakan sebuah karya berdasarkan hukum hak cipta Tiongkok. Meskipun Pengadilan mengakui bahwa isinya memenuhi persyaratan formalitas sebuah karya sastra dan pemilihan, penilaian, dan analisis data yang relevan mencerminkan orisinalitas sampai batas tertentu, Pengadilan lebih lanjut memutuskan bahwa orisinalitas bukanlah syarat yang cukup untuk memenuhi syarat sebuah karya, sedangkan orang yang menciptakan karya tersebut merupakan syarat yang diperlukan dalam hukum hak cipta Tiongkok. Lebih jauh lagi, meskipun dalam kasus *Shenzhen Tencent v Yingxun Tech* pada bulan Desember 2019 Pengadilan Distrik Shenzhen Nanshan memberikan putusannya yang mendukung karya sastra yang diproduksi AI yang dilindungi berdasarkan hukum hak cipta, Pengadilan menolak untuk mengakui bahwa karya yang dihasilkan perangkat lunak dilindungi oleh hak cipta berdasarkan alasan bahwa perangkat lunak bukanlah seorang penulis. Pengadilan dalam kedua kasus tersebut mencoba menemukan intervensi manusia dalam aktivitas kreatif.

Undang-Undang Hak Cipta, Desain, dan Paten 1988 (CDPA) di Inggris menyatakan 'penulis, dalam kaitannya dengan suatu karya, berarti orang yang menciptakannya'. Jelas, definisi ini mengecualikan ruang apa pun untuk AI sebagai penulis. Menariknya, CDPA menyediakan bagian khusus untuk karya yang dihasilkan komputer, yang mungkin paling dekat dengan situasi karya yang dibantu AI, dengan mengatakan '[d]alam kasus karya sastra, drama, musik atau seni yang dihasilkan komputer, penulis akan dianggap sebagai orang yang

melakukan pengaturan yang diperlukan untuk penciptaan karya tersebut'. Dengan kata lain, bahkan jika sebuah karya musik dibuat oleh mesin AI, pembuatnya bukanlah mesin itu sendiri tetapi orang atau orang-orang yang menyediakan basis data dan membuat algoritma untuk menginstruksikan mesin tentang cara membuat musik. Ini adalah pemrogram komputer atau insinyur perangkat lunak daripada sistem AI. Namun, tidak ada aturan yang setara untuk rekaman suara, dan hak untuk diidentifikasi sebagai penulis atau sutradara serta mempertahankan hak moral tidak berlaku untuk karya yang dihasilkan komputer. CDPA tidak memberikan perlakuan yang setara kepada penulis manusia atas karya yang dihasilkan komputer karena mereka mungkin dianggap memberikan kontribusi intelektual yang lebih kecil terhadap karya tersebut, dibandingkan dengan karya normal (misalnya karya sastra) yang sepenuhnya dibuat oleh manusia. Di Uni Eropa, baik Arahan Perangkat Lunak maupun Arahan Basis Data mengizinkan Negara Anggota UE untuk mendefinisikan kepengarangan program komputer tetapi menegaskan bahwa 'penulis basis data/program komputer adalah orang perseorangan atau sekelompok orang perseorangan yang menciptakan basis/program tersebut ...'.

Oleh karena itu, dapat dikatakan bahwa sulit untuk mengupayakan kemungkinan kepengarangan mesin dalam kerangka hak cipta yang ada (terlepas dari tingkat nasional dan internasional).

ORISINALITAS

Kriteria perlindungan hak cipta harus menilai apakah sistem legislatif menyiratkan bahwa seorang penulis harus berkontribusi pada sebuah karya yang dilindungi hak cipta. Hukum hak cipta melindungi ekspresi suatu gagasan, alih-alih gagasan itu sendiri, yang mana orisinalitas merupakan faktor kunci untuk memperoleh hak cipta. Sistem hak cipta di Inggris memiliki persyaratan orisinalitas yang rendah, yaitu suatu karya dianggap orisinal selama karya tersebut mencerminkan keterampilan, upaya, dan penilaian yang memadai dari sang penulis. Meskipun ada kemungkinan perubahan segera pada persyaratan orisinalitas di Inggris setelah kasus *Infopaq International A/S v Danske Dagbaldes Forening*, tidak diragukan lagi bahwa kata-kata "keterampilan" dan "tenaga kerja" hanya dapat dikaitkan dengan atau menghubungkan ke manusia. Selain itu, baik Arahan Perangkat Lunak dan Arahan Basis Data di UE menyatakan bahwa orisinalitas berarti kreasi intelektual seorang penulis sendiri. Kasus *Infopaq* lebih lanjut menyoroti hubungan antara orisinalitas dan kreasi intelektual penulis sendiri, dan interpretasinya tentang orisinalitas dipahami sebagai cerminan dari kepribadian seorang penulis sendiri. 'Kepribadian mengacu pada perbedaan individu dalam pola karakteristik berpikir, merasa dan berperilaku'. Tidak ada keraguan bahwa kepribadian memperjelas perlunya manusia sebagai penulis suatu karya, dan ini juga menunjukkan kebebasan penulis untuk membuat pilihan kreatif. Manusia memiliki kebebasan penuh selama kegiatan kreatif, meskipun praktik mereka dapat dibatasi oleh teknologi, nilai, dan sistem sosial pada saat itu. Ini mungkin alasan mengapa CJEU dalam kasus *Football Dataco Ltd and Others v Yahoo! UK Ltd and Others* menekankan ketidakmungkinan membentuk pilihan bebas dan kreatif jika membuat sebuah karya adalah hasil dari 'pertimbangan teknis, aturan, atau batasan'.

Dalam hal ini, karya musik yang dihasilkan oleh mesin tidak mungkin, sehingga memenuhi persyaratan orisinalitas dalam konteks hak cipta. Dengan pengakuan penggunaan AI dan aktivitas pembelajaran mesin terkaitnya dan pemahaman bahwa kerangka hak cipta saat ini menghalangi kemungkinan AI diberi status pencipta, Uni Eropa telah berupaya untuk menangani masalah hukum terkait robotika dalam Rekomendasi 2017 kepada Komisi Aturan Hukum Perdata tentang Robotika. Namun, hal ini masih jauh dari mengklarifikasi masalah apakah sistem AI memiliki karya yang 'diciptakannya'.

7.5 SOLUSI ALTERNATIF KARYA YANG DIHASILKAN KOMPUTER

Meskipun sistem hak cipta nasional dan internasional hanya memberikan sedikit ruang untuk menafsirkan kepengarangan non-manusia dan orisinalitas, tetap penting untuk mengeksplorasi solusi alternatif guna menangani karya yang dihasilkan AI dalam konteks hak cipta. Hal ini berharga tidak hanya untuk merangsang kreasi terkait AI dan untuk tujuan investasi, tetapi juga untuk mengatasi ketidakpastian dan masalah hak cipta yang telah diperkirakan.

Penting untuk membahas solusi alternatif yang dapat diterapkan dalam sistem hak cipta saat ini. Hukum hak cipta Inggris menyatakan bahwa pencipta karya yang dihasilkan komputer adalah 'orang yang melakukan pengaturan yang diperlukan untuk penciptaan karya tersebut'. Hal ini tampaknya menjadi jalan yang memungkinkan bagi musik yang dihasilkan AI untuk mendapatkan perlindungan hak cipta. Namun, ada pandangan yang meragukan arti pasti dari kata 'pengaturan' karena menyiratkan rencana dan persiapan untuk mewujudkan sesuatu (di sini berarti merencanakan dan mempersiapkan untuk menciptakan musik), dan orang-orang yang membuat pengaturan tersebut beragam; mereka bisa jadi pengguna, perancang program, investor perangkat lunak, atau instruktur yang melatih dan menginstruksikan para programmer. Dalam *Nova Productions Ltd v Mazooma Games Ltd & Ors*, Pengadilan memutuskan bahwa perancang permainan adalah orang yang membuat pengaturan, bukan pengguna yang memainkan permainan dan menghasilkan gambar unik selama bermain. Namun, ini tidak mudah untuk pembuatan musik. Sementara musisi profesional mengadopsi perangkat lunak musik AI selama proses kreatif, mereka lebih suka melihat perangkat lunak AI sebagai alat yang berguna daripada mendominasi kreasi mereka.

Namun, faktanya, musisi mungkin secara tidak sadar terinspirasi dan dipengaruhi oleh generasi otonom AI. Seluruh proses kreatif adalah uji coba menurut penyesuaian, penilaian, dan pemilihan musisi atas sumber-sumber relevan yang disediakan oleh perangkat lunak AI, yang mencakup elemen dampak-kolaborasi interaktif: musisi terinspirasi oleh uji coba musik yang diproduksi AI, dibuat berdasarkan penyesuaian, penilaian, dan pemilihan sumber mereka sendiri; perangkat lunak AI terus-menerus menghasilkan musik yang sesuai dengan elemen dan pilihan tambahan yang secara bertahap mencocokkan musik dengan ide-ide kreatif musisi. Musisi dapat memodifikasi, mengadaptasi, dan mengintegrasikan percobaan ini dengan ide mereka sendiri, dan sebuah karya akhir pun tercipta. Dari sudut pandang penulis, karya musik yang dihasilkan dengan cara ini juga harus diperlakukan sebagai musik yang dibantu AI karena AI memainkan peran asisten yang penting dalam membantu musisi manusia menyelesaikan

karya akhir. Dari perspektif ini, rasanya aneh untuk mengatakan bahwa perancang/insinyur perangkat lunak musik AI adalah rekan pencipta, dan yang lebih penting, musik yang dihasilkan AI menghasilkan batas yang tidak jelas antara musik yang diciptakan manusia dan musik yang dihasilkan komputer. Sulit untuk mengidentifikasi bagian mana yang merupakan ciptaan intelektual musisi manusia dan bagian mana yang merupakan hasil upaya sistem AI. Dari perspektif penulis, karya ini lebih seperti karya kolaboratif dengan status rekan penulis. Terlebih lagi, kita dapat memperkirakan perangkat lunak AI di masa depan mampu mengambil keputusan atau mencapai pengiring yang sebagian besar otonom sambil melibatkan proses kreasi musisi manusia. Dari sudut pandang ini, partisipasi AI dalam kreasi melampaui cakupan alat semata dalam karya yang dihasilkan komputer tradisional, sehingga kategori karya yang dihasilkan komputer tidak akan cukup terdefinisi untuk jenis musik yang diproduksi AI tersebut (baik musik yang dibantu AI maupun yang dihasilkan). Meskipun demikian, sistem hak cipta menyatakan dengan jelas bahwa hanya penulis manusia yang memenuhi syarat untuk hak cipta.

Karya yang dibuat untuk disewa atau karya badan hukum?

Beberapa akademisi membahas kemungkinan memperlakukan karya yang dihasilkan AI sebagai karya yang dibuat untuk disewa yang diatur dalam hukum hak cipta AS, dan isu yang umum adalah 'bagaimana mesin memiliki kecerdasan yang cukup untuk mengambil keputusan menyetujui atau tidak menyetujui perjanjian kontraktual?' Lebih lanjut, terdapat kekhawatiran lebih lanjut tentang bagaimana AI, perancang programnya, dan penggunaanya mengakomodasi hubungan komisioner-pencipta atau hubungan pemberi kerja-karyawan dalam karya yang dihasilkan AI. Lagipula, tidaklah lazim memperlakukan AI sebagai karyawan tanpa menjadikannya 'dalam proses kerja'. Selain itu, keluaran yang dihasilkan AI tidak termasuk dalam situasi pekerjaan komisi karena tidak ada bukti yang menunjukkan bahwa AI dititipkan oleh manusia/perusahaan untuk menyelesaikan pekerjaan komisi.

Kasus *Tencent Shenzhen* di Tiongkok memberi sedikit harapan bahwa karya yang dihasilkan AI dapat dipandang sebagai badan hukum – sebuah hadiah Natal yang indah untuk karya yang dihasilkan AI (putusan dibuat pada 24 Desember 2019). Perangkat lunak yang dikembangkan oleh Tencent (Penggugat) menyebut Dreamwriter sebagai sistem bantuan penulisan cerdas yang telah membantu mereka menyelesaikan 300.000 karya per tahun sejak 2015. Sebuah artikel keuangan oleh Dreamwriter dipublikasikan di situs web Tencent Stock pada 20 Agustus 2018, dan kemudian disalin secara lengkap dan tersedia untuk umum melalui situs web tergugat. Pengadilan memutuskan bahwa artikel tersebut merupakan karya sastra dari perspektif formalitas – mampu direproduksi dan menjadi dasar kemunculannya. Kunci untuk membentuk sebuah karya sastra dalam kasus ini adalah persyaratan orisinalitas. Pengadilan menyatakan bahwa konten artikel tersebut mencerminkan pilihan, analisis, dan penilaian data dan informasi saham pada hari itu, dengan logika dan ekspresi yang jelas serta struktur yang wajar – oleh karena itu, artikel tersebut orisinal.

Ketika menentukan apakah artikel ini mencerminkan pilihan, penilaian, dan keterampilan individu pencipta, Pengadilan mencatat perbedaan antara penciptaan artikel ini dan metode umum penciptaan sebuah artikel, dan mengakui adanya kesenjangan waktu

antara pilihan dan pengaturan kelompok pencipta dan tindakan penulisan aktual yang dieksekusi oleh perangkat lunak Dreamwriter. Namun, Pengadilan percaya bahwa kurangnya sinkronisasi dihasilkan dari karakteristik proses teknis dan menolak untuk memperlakukan proses eksekusi otomatis sebagai penciptaan karena hal ini didasarkan pada memperlakukan perangkat lunak sebagai pencipta sampai batas tertentu. Menariknya, Pengadilan menyatakan bahwa penciptaan artikel ini secara langsung dihasilkan dari aktivitas karyawan Tencent dalam pengaturan dan pemilihan individu, dan perangkat lunak hanya menyediakan alat teknis. Pengadilan enggan menerima operasi otonom AI sebagai proses penciptaan intelektual, tetapi sebaliknya memperlakukannya hanya sebagai alat pintar dan lebih suka melacak intervensi manusia sebanyak mungkin. Pengadilan memperlakukan artikel ini sebagai penciptaan intelektual terintegrasi yang dibuat oleh penilaian dan pilihan karyawan dan pengoperasian perangkat lunak - dan bahwa artikel tersebut adalah karya sastra yang dilindungi oleh hukum hak cipta Tiongkok. Lebih jauh lagi, menurut pasal 11 undang-undang hak cipta di Tiongkok, '[d]i mana suatu karya diciptakan menurut maksud dan di bawah pengawasan dan tanggung jawab suatu badan hukum atau organisasi lain, badan hukum atau organisasi tersebut akan dianggap sebagai pencipta karya tersebut'. Artikel yang disengketakan adalah ciptaan intelektual terpadu yang dibuat oleh beberapa tim yang berbagi tugas yang berbeda, yang mencerminkan kebutuhan Tencent untuk menerbitkan artikel keuangan. Seluruh kelompok karyawan yang membuat artikel tersebut memiliki beberapa tim yang diawasi Tencent; artikel ini diterbitkan di situs web Tencent dengan catatan di akhir yang menyatakan, 'artikel ini secara otomatis ditulis oleh robot Tencent Dreamwriter', yang menunjuk ke Tencent sebagai pencipta (penggugat). Pengadilan menyatakan bahwa ini menunjukkan Tencent akan menanggung semua kewajiban relevan yang berasal dari artikel yang disengketakan. Oleh karena itu, Pengadilan memutuskan bahwa artikel yang disengketakan adalah karya suatu badan hukum, dan Tencent, penggugat, dianggap sebagai pencipta dan karenanya berhak atas hak cipta.

Dalam kasus ini, Pengadilan menekankan hubungan langsung antara bentuk ekspresi dalam artikel yang disengketakan dan aktivitas intelektual kelompok karyawan dalam mengatur dan memilih masukan, menetapkan ketentuan, templat, dan gaya bahasa. Sampai batas tertentu, Pengadilan lebih berpihak pada pengguna (kelompok karyawan) dan mengecilkan peran otonom perangkat lunak. Mengingat perancang program dan pengguna program semuanya berasal dari Tencent, tidak ada perselisihan mengenai kepengarangan Tencent. Namun, jika pengguna dan perancang program merupakan pihak yang independen dan terpisah, putusan ini dapat dianggap bias terhadap pengguna program, yang bertentangan dengan preferensi perancang gim dalam kasus *Nova*. Sayangnya, Pengadilan dalam kasus *Tencent di Shenzhen* tidak membuat hubungan langsung antara orisinalitas artikel dan perancang perangkat lunak. Oleh karena itu, permasalahannya masih belum jelas. Model Jukedek tampaknya menjawab pertanyaan ini sampai batas tertentu: model ini memungkinkan penggunaannya untuk menghasilkan lima lagu gratis sebulan sebelum mereka mulai membayar untuk setiap lagu, tetapi pengguna tetap harus membayar lebih untuk mendapatkan hak cipta. Dengan kata lain, perancang/pemilik program memiliki hak cipta,

sementara penggunaannya yang menciptakan lagu hanya menikmati penggunaan pribadi yang terbatas.

Sistem hak cipta terutama dirancang untuk mengatur aktivitas kreatif manusia. Dari sudut pandang penulis, tidaklah tepat untuk terlalu terlibat dalam mencari tenaga intelektual manusia dalam barang-barang yang diproduksi oleh teknologi, karena ketidaktahuan akan kontribusi teknologi dapat menghambat inovasi teknologi dan penyebaran produksi budaya terkait. Lagipula, hal ini dapat memprediksi bahwa penelitian AI dan pengembangannya menuju masa depan yang sangat otonom dan pengambilan keputusan mandiri, yang sama sekali berbeda dari teknologi-teknologi yang masih sangat bergantung pada pengambilan keputusan manusia. Oleh karena itu, mengakui kontribusi teknologi terhadap penciptaan dan arah masa depannya, merangkul musik yang diproduksi oleh AI dengan intervensi manusia yang minimal atau nol, adalah kunci untuk mempertimbangkan sistem hukum yang tepat guna merespons perubahan radikal yang akan datang di sektor musik. Satu hal yang harus kita sadari adalah bahwa satu undang-undang dapat menyelesaikan beberapa masalah tetapi tidak semuanya.

7.6 KESIMPULAN

Meskipun teknologi berkembang lebih cepat dari sebelumnya dan terus menciptakan hal-hal baru dengan sifat/fitur yang sangat berbeda, undang-undang yang digerakkan oleh manusia selalu tertinggal dan tertantang oleh isu-isu kompleks ini. Kemunculan sistem AI dapat menjungkirbalikkan atau setidaknya menantang serangkaian pandangan tradisional mengenai pengarang, kreativitas, dan proses penciptaan, baik dalam konteks akal sehat manusia maupun hak cipta. Meskipun manusia dan sistem AI memulai hubungan kompetitif dalam industri musik, memahami sifat sistem AI dan musik yang diproduksi AI, serta mempertanyakan status hukum AI dalam sistem hak cipta yang ada, sangatlah penting untuk mengidentifikasi isu-isu yang muncul, seperti pelanggaran, remunerasi, dan tanggung jawab hukum pihak-pihak terkait. Dapat dilihat bahwa sistem hak cipta yang ada berusaha memaksimalkan perlindungannya terhadap perubahan teknologi, sehingga terdapat berbagai hak *sui generis* di bawah payung sistem hak cipta yang telah muncul, tetapi tidak satu pun dari hak-hak tersebut yang menyentuh dasar-dasar kerangka kerjanya. Namun, kali ini, sistem hak cipta yang ada tampaknya macet, karena revolusi radikal sistem AI menantang doktrin fundamentalnya: penciptaan oleh manusia dan kreativitas/orisinalitas manusia. Meskipun demikian, ada baiknya mempertimbangkan kemungkinan menafsirkan ulang istilah dan doktrin kunci dalam konteks hak cipta untuk memetakan perubahan teknologi ini. Mungkin perlindungan cahaya untuk musik yang dibantu AI di bawah sistem hak cipta yang ada dapat dipertimbangkan. Namun, sejauh mana perlindungan cahaya ini seharusnya? Untuk musik yang dihasilkan AI, cara terbaik, dari sudut pandang penulis, adalah membangun sistem baru yang khusus untuk robotika untuk menyelesaikan masalah yang relevan, tetapi bagaimana menyelaraskan sistem baru ini dengan hukum lain yang mengatur aktivitas manusia juga merupakan pertanyaan besar.

Faktanya, mengidentifikasi sistem AI dan hak cipta karya mereka hanyalah titik awal dari masalah hak cipta terkait AI. Ada banyak masalah di sekitarnya yang juga menantang atau akan menantang seluruh industri musik. Salah satu masalah yang menjadi perhatian, yang diangkat oleh Percakapan WIPO tentang AI dan IP, adalah kemungkinan pelanggaran data pelatihan AI. Alasan sistem AI dapat 'menciptakan' musik adalah karena mereka mengandalkan basis data yang kuat untuk mendukung pembelajaran mesin mereka, yang mereka gunakan untuk menghasilkan karya musik. Oleh karena itu, apakah dan bagaimana melindungi atau/dan mengatur basis data adalah pertanyaan besar. Namun, satu masalah di sini adalah bahwa informasi pada basis data pelatihan terdiri dari musik atau produk yang dibuat oleh manusia. Misalnya, seorang komposer AI perlu mempelajari sejumlah besar musik yang dibuat manusia untuk mengumpulkan fitur-fiturnya dan menentukan aturan sebelum menghasilkan karya musik apa pun. Jika karya musik yang dibuat manusia ini dapat dilindungi hak cipta dan masih dalam jangka waktu yang dilindungi hak cipta, apakah penggunaan karya-karya ini tanpa izin dari pemilik hak cipta akan merupakan pelanggaran hak cipta adalah pertanyaan yang perlu diselidiki lebih lanjut. Lebih lanjut, sementara sistem hak cipta yang ada memberikan pembelaan – penggunaan yang wajar – tidak jelas apakah penggunaan data untuk pembelajaran mandiri AI dapat dikualifikasikan sebagai penggunaan yang wajar dalam studi dan penelitian pribadi. Selain itu, meskipun seorang komposer AI menghasilkan karya musik yang merangkul ciri-ciri musik rakyat milik kelompok etnis minoritas/komunitas adat tertentu, atau jika seorang komposer AI menghasilkan karya musik yang mengintegrasikan berbagai sumber musik rakyat, masih belum jelas apakah hal ini akan dianggap sebagai penyalahgunaan budaya musik tradisional. Tanpa identifikasi yang jelas tentang status hukum sistem AI dalam konteks hak cipta, pertanyaan-pertanyaan ini tetap sulit dijawab.

BAB 8

AI DALAM MANAJEMEN RISIKO PENERBANGAN

8.1 PENDAHULUAN

Dengan fondasi yang diletakkan oleh para pionir dan peneliti, penerbangan selalu menjadi yang terdepan dalam kemajuan. Berorientasi pada keselamatan, komunitas penerbangan memusatkan sebagian besar upayanya untuk meningkatkan perlindungan jiwa dan kesehatan mereka yang terlibat dalam operasi udara, dan untuk itu, mereka bersemangat untuk mengeksplorasi potensi teknologi yang sedang berkembang.

Pendekatan 'pelajaran yang dipetik' merupakan dasar bagi sebagian besar peningkatan keselamatan di sektor ini. Meskipun teknologi dan pengetahuan manusia telah berkembang pesat selama bertahun-tahun, industri ini tidak pernah lebih rentan daripada saat ini. Spesifikasi keuangannya (misalnya, persediaan yang mudah rusak, biaya tetap yang tinggi, kelebihan kapasitas, ketergantungan pada faktor dan siklus eksternal yang memengaruhi permintaan) ditambah dengan meningkatnya nilai dan kompleksitas pesawat modern, transformasi digital dan meningkatnya saling ketergantungan antar pelaku penerbangan, serta pandemi global, semuanya menambah beratnya tantangan yang akan datang.

Oleh karena itu, selain belajar dari peristiwa dramatis atau disruptif yang telah terjadi, para profesional penerbangan berupaya meningkatkan kapasitas mereka untuk mengambil tindakan korektif terlebih dahulu. Kepatuhan terhadap pedoman dan standar industri merupakan hal yang sangat penting dalam hal ini, demikian pula manajemen risiko yang tepat dalam arti yang lebih luas. Namun, perlu dicatat bahwa tingkat kemajuan peralatan penerbangan, dikombinasikan dengan volume operasi udara yang terjadi di dunia kontemporer, telah secara serius melemahkan kemampuan industri untuk mengenali dan mencegah bahaya tepat waktu.

Dalam konteks di atas, penggunaan AI memberikan sejumlah peluang dan memperkuat upaya manusia dalam domain transportasi udara. Bab ini membahas penerapan AI terkini dalam penerbangan. Bab ini mengkaji potensi AI untuk peningkatan yang diperlukan di bidang manajemen risiko dan menyerukan inisiatif regulasi di seluruh dunia untuk memastikan implementasi yang aman dan pengembangan AI lebih lanjut di sektor penerbangan. Bab ini juga mencakup dampak COVID-19 terhadap industri.

8.2 KEADAAN AI SAAT INI DALAM PENERBANGAN

Selama beberapa dekade, sektor penerbangan telah memperdebatkan kegunaan AI, tetapi baru dalam beberapa tahun terakhir potensi AI telah memicu minat serius di antara para pelaku industri. Beberapa maskapai penerbangan, bandara, dan produsen bahkan telah mengamati teknologi AI lebih dekat dan memutuskan untuk berinvestasi dalam pengembangan solusi AI. Pergeseran ini sebagian dapat dikaitkan dengan peningkatan signifikan dalam kemampuan teknologi, termasuk perkembangan di berbagai bidang seperti ketersediaan data, daya komputasi, kapasitas penyimpanan, dan, sebagai hasilnya,

aksesibilitas dan keterjangkauan AI untuk bisnis yang terkait dengan penerbangan. Hingga saat ini, penggunaan AI dalam penerbangan sangat terbatas, tetapi ketidakmatangan implementasinya tidak boleh menjadi alasan untuk mengabaikan relevansi dan potensi dampak AI terhadap sektor ini.

Dalam hal keberadaan AI saat ini dalam penerbangan, pemeriksaan laporan industri, peta jalan, dan kebijakan mengungkapkan bahwa penekanan awal terutama pada manajemen yang efektif dari arus penumpang dan pesawat yang besar. Banyak pemangku kepentingan penerbangan memutuskan untuk mengeksplorasi berbagai solusi berbasis AI yang ditujukan untuk mengatasi masalah kapasitas, seperti masalah penundaan dan kepuasan penumpang. 'Agen virtual dan chatbot' dan analitik prediktif telah menjadi dua aplikasi AI yang paling umum. Yang pertama telah mengotomatiskan dan mempercepat komunikasi pelanggan, sedangkan yang terakhir telah memungkinkan, misalnya, peramalan permintaan untuk produk tertentu atau kemungkinan kegagalan peralatan sebelum terjadi. Selain itu, potensi aplikasi AI di bidang penjadwalan otomatis untuk penerbangan dan kru, dan manajemen lalu lintas udara telah dipertimbangkan.

Meskipun penggunaan AI dalam bisnis penerbangan komersial saat ini berada di luar perhatian publik, industri ini telah memulai beberapa proyek yang bertujuan untuk mengatasi masalah ini secara komprehensif. Misalnya, pada tahun 2018, Asosiasi Transportasi Udara Internasional (IATA) menerbitkan 'Buku Putih AI dalam Penerbangan', dengan tujuan untuk meningkatkan pengetahuan tentang AI di seluruh sektor penerbangan. Tujuan lainnya termasuk memberikan penjelasan yang lebih luas tentang manfaat AI bagi industri, mengidentifikasi risiko dan prospek yang terkait dengan penggunaan AI, dan memberikan rekomendasi mengenai dimulainya penerapan AI dalam bisnis penerbangan. Dokumen ini menjelaskan dasar-dasar AI, menyajikan kapabilitas dan studi kasusnya untuk kepentingan sektor penerbangan. Dokumen ini memberikan informasi tentang inisiatif IATA di bidang ini dan menyatakan bahwa AI pasti akan hadir di sektor ini. Menurut IATA, industri dapat mengabaikan kebenaran ini dengan risikonya sendiri atau mengeksplorasi kemungkinan AI dan secara aktif terlibat dalam revolusi AI. Dokumen ini lebih lanjut menguraikan bahwa meskipun banyak tantangan yang harus diatasi, merangkul AI lebih cepat daripada yang lain adalah satu-satunya cara untuk tetap menjalankan bisnis dalam jangka panjang.

Baru-baru ini, berbagai kegiatan terkait AI telah dilakukan di sektor penerbangan Eropa. Misalnya, pada tahun 2020 Badan Keselamatan Penerbangan Uni Eropa (EASA) mengembangkan peta jalan untuk AI dalam penerbangan. Tujuan utama dari dokumen ini adalah untuk membahas peran AI dalam penerbangan, untuk menentukan tujuan yang harus dicapai dan langkah-langkah yang harus diikuti untuk menanggapi pertanyaan yang muncul terkait dengan kepercayaan publik, etika, sertifikasi dan proses standardisasi. Juga, pada tahun 2020, European Aviation/ATM AI High Level Group dibentuk. Dipimpin oleh Organisasi Eropa untuk Keselamatan Navigasi Udara (EUROCONTROL), Grup tersebut menyatukan kelas perwakilan yang luas, termasuk maskapai penerbangan, bandara, penyedia layanan navigasi udara (ANSP), produsen, Komisi Eropa, serta organisasi militer dan staf, untuk secara komprehensif membahas masalah AI dalam laporan mereka. Melalui dokumen ini, para

peserta yang terlibat bermaksud untuk meningkatkan pemahaman tentang AI di sektor penerbangan dan menunjukkan konsekuensi utamanya bagi entitas swasta dan publik, untuk persyaratan sipil atau militer. Bersamaan dengan itu, mereka menyatakan komitmen untuk mengembangkan strategi AI mereka sendiri yang relevan. Lebih lanjut, mereka berupaya mengidentifikasi hambatan terhadap penggunaan AI yang lebih luas dalam spektrum penuh operasi penerbangan/ATM dan menentukan modifikasi kerangka kerja yang diperlukan untuk mendorong pengembangan AI dalam penerbangan Eropa.

COVID-19 tidak diragukan lagi berdampak pada pertimbangan dan tujuan ini. Industri ini, yang berfokus pada keberlangsungan hidup dan tantangan dalam mengurangi pengeluaran kasnya, menempatkan nilai terbesar pada fleksibilitas dan ketahanan. Langkah selanjutnya dan cara paling pasti untuk pulih dari krisis ini adalah dengan menerima kemungkinan bahwa COVID-19 akan menjadi endemik, dan kehadirannya akan berdampak untuk waktu yang lama. Dengan realisasi hal ini muncul pemahaman bahwa baik pemerintah maupun industri tidak dapat terus-menerus menjalankan strategi penghindaran risiko sepenuhnya tanpa batas waktu. Risiko COVID-19 perlu dikelola secara efektif, pesawat perlu terbang, dan perbatasan perlu dibuka. Hal ini, pada gilirannya, menuntut tindakan cepat dan inisiatif regulasi untuk memungkinkan akses praktis dan penerapan solusi mitigasi risiko yang tersedia.³⁰ Seiring dengan tersedianya bukti ilmiah yang kuat, mekanisme manajemen risiko yang disebabkan oleh COVID-19 tersebut harus terus-menerus dinilai ulang dan diganti dengan metode yang terbukti lebih efektif.

Dalam hal ini, kebutuhan untuk memulihkan, memulihkan kepercayaan, dan memenuhi kekhawatiran publik terkait perlindungan kesehatan penumpang dan profesional penerbangan kemungkinan akan mendorong kemajuan jangka panjang di bidang AI sebagai alat manajemen risiko dan peningkatan keselamatan dalam penerbangan. Misalnya, AI dapat mendukung pengendalian infeksi secara real-time, membantu penyaringan penumpang, dan meningkatkan pelacakan penyebaran virus serta identifikasi kasus berisiko tinggi. Selain itu, berkat aplikasi AI seperti pemeliharaan prediktif, AI dapat mendukung industri penerbangan dalam hal keselamatan operasional. Khususnya, sangat penting bagi penerbangan komersial untuk terus membuktikan bahwa perjalanan udara dapat beroperasi dengan aman dengan mematuhi aturan dan proses untuk tingkat keselamatan tertinggi. Pada saat yang sama, perubahan signifikan, yang melibatkan keselamatan kesehatan, dampak lingkungan, dan penerapan teknologi baru, dalam bisnis perlu diterima dan dipahami untuk memastikan bahwa industri penerbangan kembali lebih kuat dan mampu mengatasi tantangan baru yang pasti akan muncul.

COVID-19 telah menghentikan hampir segalanya, termasuk bisnis penerbangan yang sederhana, tetapi indikasi awal menunjukkan bahwa industri penerbangan yang akan muncul dari krisis COVID-19 akan lebih bersemangat untuk menerapkan solusi berbasis AI. Namun kali ini, fokusnya telah berubah dan sebagian besar upaya diharapkan bergeser dari peningkatan pengalaman pelanggan atau pengelolaan masalah kapasitas menjadi peningkatan keselamatan kesehatan dan pembatasan interaksi antarmanusia menggunakan berbagai sarana perjalanan nirsentuh.

8.3 PERAN AI DALAM PENINGKATAN KESELAMATAN PENERBANGAN DI INDONESIA

Pemanfaatan kecerdasan buatan dalam industri penerbangan di Indonesia memiliki peranan yang sangat penting dalam manajemen potensi risiko yang dapat mengganggu keselamatan penerbangan. Penerapan teknologi ini dilakukan dengan mengintegrasikan sistem analisis data berbasis AI pada berbagai aspek operasional, mulai dari pemeliharaan pesawat, pengelolaan lalu lintas udara, hingga pengambilan keputusan oleh pilot dan maskapai. Salah satu contoh penerapan yang paling nyata adalah melalui sistem predictive maintenance. Dengan memanfaatkan data yang dikumpulkan oleh ribuan sensor pada pesawat, AI mampu menganalisis kondisi komponen secara real-time dan mendeteksi adanya pola abnormal yang tidak mudah diidentifikasi secara manual. Deteksi dini ini memungkinkan pihak maskapai untuk melakukan tindakan pencegahan sebelum kerusakan benar-benar terjadi, sehingga risiko gangguan penerbangan maupun kecelakaan dapat diminimalisasi secara signifikan. Selain dalam aspek pemeliharaan, AI juga memiliki pengaruh besar terhadap mitigasi risiko yang bersumber dari faktor eksternal, khususnya kondisi cuaca. Indonesia sebagai negara kepulauan dengan iklim tropis memiliki tingkat dinamika cuaca yang tinggi, termasuk potensi badai, hujan lebat, dan angin kencang yang kerap memengaruhi keselamatan penerbangan. Melalui sistem analisis prediktif berbasis AI, data cuaca dapat diproses lebih cepat dan akurat untuk memberikan peringatan dini kepada pilot maupun pengatur lalu lintas udara. Dengan demikian, keputusan untuk mengubah jalur penerbangan atau menunda keberangkatan dapat dilakukan secara lebih tepat waktu sehingga mengurangi potensi risiko yang diakibatkan oleh faktor alam.

AI juga mulai diterapkan dalam sistem manajemen lalu lintas udara (*air traffic management*), di mana teknologi ini digunakan untuk mengolah data penerbangan yang kompleks guna mengoptimalkan jalur pesawat di udara. Hal ini sangat relevan di Indonesia yang memiliki ruang udara padat pada jalur tertentu, seperti rute Jakarta–Denpasar atau Jakarta–Makassar. Dengan adanya AI, risiko tabrakan di udara maupun keterlambatan akibat kepadatan lalu lintas dapat ditekan melalui pengaturan rute yang lebih efisien dan terukur. Lebih jauh lagi, AI berkontribusi dalam mendukung implementasi Safety Management System (SMS) yang diwajibkan oleh International Civil Aviation Organization (ICAO) dan diadopsi dalam regulasi Kementerian Perhubungan Indonesia, dengan cara meningkatkan aspek predictive safety sehingga manajemen risiko lebih bersifat proaktif dibandingkan reaktif.

Contoh nyata penerapan teknologi ini di Indonesia dapat dilihat dari langkah Garuda Indonesia yang mulai mengadopsi sistem berbasis data dan kecerdasan buatan untuk memantau kesehatan armadanya. Melalui kerja sama dengan perusahaan teknologi penerbangan internasional, maskapai ini dapat melakukan analisis prediktif terhadap kondisi mesin pesawat sehingga perawatan tidak hanya dilakukan berdasarkan jadwal rutin, tetapi juga berdasarkan data aktual yang menunjukkan kebutuhan perawatan. Sementara itu, AirNav Indonesia telah memanfaatkan teknologi digital berbasis AI untuk mendukung pengelolaan lalu lintas udara, termasuk dalam pemetaan jalur penerbangan serta analisis cuaca. Sistem ini membantu mengantisipasi risiko penundaan dan gangguan penerbangan, terutama pada bandara-bandara dengan tingkat kepadatan tinggi seperti Soekarno-Hatta di Jakarta.

Dengan berbagai penerapan tersebut, dapat dipahami bahwa AI tidak hanya berfungsi sebagai alat bantu teknologi, melainkan juga sebagai instrumen strategis dalam membangun sistem manajemen risiko penerbangan yang lebih modern, adaptif, dan selaras dengan standar keselamatan internasional. Kehadiran AI mendorong terjadinya pergeseran paradigma dalam manajemen risiko penerbangan di Indonesia, yaitu dari pendekatan tradisional yang lebih mengandalkan inspeksi manual dan respons terhadap insiden, menuju pendekatan prediktif yang berbasis pada data dan analisis otomatis. Hal ini menunjukkan bahwa perkembangan teknologi AI sangat relevan untuk meningkatkan keselamatan penerbangan nasional, sekaligus menjadi tantangan bagi pemangku kepentingan untuk terus memperbarui kebijakan dan infrastruktur agar dapat memaksimalkan manfaat yang ditawarkan oleh teknologi tersebut.

8.4 AI DAN POTENSI MANAJEMEN RISIKO PENERBANGAN

Paradoksnya, dari perspektif manajemen risiko, risiko yang sepenuhnya baru bukanlah perhatian utama. Namun, kesulitan utamanya terletak pada risiko yang sudah ada tetapi menjadi lebih sulit dideteksi, atau yang mulai muncul dengan cara yang tidak diketahui dan khusus. Contoh terbaru adalah risiko yang terkait dengan kesehatan dan kebugaran penumpang untuk terbang. Peluang untuk menyediakan para profesional penerbangan dengan keahlian dan pelatihan yang memadai untuk mendeteksi dan mengelola beragam potensi ancaman yang dibawa ke dalam dan di sekitar pesawat sangatlah kecil. Kompleksitas dan dinamika risiko akan selalu menantang upaya manusia dalam hal ini.

Bagi mereka yang terlibat dalam aktivitas penerbangan, terwujudnya risiko dapat berarti kehilangan atau kerusakan properti atau aset, kematian atau cedera staf penerbangan, atau kerugian dan kerusakan yang mengakibatkan kewajiban kepada pihak ketiga. Untuk lebih memahami konteks dan fluktuasi risiko yang terus berubah, penerbangan membutuhkan alat yang lebih baik untuk mengenali dan memahami anomali yang terjadi. Dalam hal ini, AI menawarkan berbagai peluang.

Dalam hal kesiapsiagaan risiko, apa yang tampak sangat menjanjikan adalah kapasitas AI untuk belajar, mengembangkan diri, dan, tanpa penundaan, menerapkan informasi yang diperoleh pada situasi atau masalah yang belum pernah dialami sebelumnya. Dalam hal ini, AI mampu menangkap tren yang tidak dapat dikenali oleh sistem analitik manusia maupun tradisional. AI menyerap sejumlah besar data baik secara historis maupun waktu nyata, dan menggunakannya untuk menangkap perbedaan tersebut. Fitur ini sangat penting, karena dalam kebanyakan kasus, ancaman yang muncul bersamaan dengan ketidakaturan pola. Selain itu, hal ini membuat AI sangat cocok untuk industri penerbangan, yang sangat bergantung pada aliran data digital antara sistem udara dan darat.

Pertumbuhan jumlah informasi yang dipertukarkan secara stabil menjadikan sektor penerbangan sangat kaya akan data. Namun, dalam hal pemrosesan data, perangkat yang digunakan saat ini tidak mampu mengelola volume informasi yang dikumpulkan dari sistem pesawat, pengawasan, maskapai penerbangan, bandara, atau pemangku kepentingan industri lainnya. Sebagian karena sekitar 90 persen data yang dihasilkan tidak terstruktur, yang berarti tidak dapat diposisikan dan dianalisis dengan cara konvensional dalam kolom dan baris. Dalam

konteks ini, keunggulan AI terletak pada kemampuannya untuk menangani data yang tidak terstruktur. Saat ini, kemampuan kognitif, termasuk penambahan data, pembelajaran mesin, dan pemrosesan bahasa alami, dapat mengambil alih metode analisis lama dan menarik kesimpulan dari data yang tidak terstruktur, sehingga menghasilkan deteksi bahaya yang diketahui dan tidak diketahui yang lebih baik dan lebih cepat.

Dalam konteks manajemen risiko, keunggulan utama AI adalah: AI dapat dibagi menjadi beberapa area berikut:

- 1) Pemrosesan data: penggunaan data terstruktur dan tidak terstruktur dalam jumlah besar; pembaruan pola dan variasi kumpulan data;
- 2) Efisiensi: peningkatan otomatisasi dukungan rutin dalam proses manajemen risiko;
- 3) Waktu nyata dan prediktif: identifikasi risiko yang muncul, penerbitan rekomendasi mitigasi risiko, reaksi yang lebih cepat dan akurat terhadap situasi yang berpotensi membahayakan;
- 4) Pengambilan keputusan bisnis: proses yang lebih baik berkat pemahaman dan pengenalan risiko yang lebih baik.

Secara khusus, dalam bisnis penerbangan, AI dapat mendukung proses manajemen risiko dengan memungkinkan pemantauan keselamatan pesawat secara real-time dan pemantauan kesehatan penumpang, meningkatkan perlindungan aset, meningkatkan alokasi sumber daya dinamis (misalnya, meminimalkan antrean di bandara melalui analitik prediktif dan aplikasi notifikasi penumpang), dan seterusnya.

Potensi AI membangkitkan imajinasi industri transportasi udara, memicu visi masa depan di mana para pelaku penerbangan umumnya mengadopsi AI untuk berbagai tujuan, termasuk penerapan AI di lingkungan kokpit. Meskipun ketika mengincar transisi AI, tidak masalah untuk memiliki tujuan yang ambisius, penting juga untuk memulai dengan bukti konsep. Kegunaan dan kelayakan komersial dari tujuan tersebut untuk entitas tertentu, seperti maskapai penerbangan, bandara, atau organisasi penerbangan, dapat dibuktikan melalui pengujian atau proyek percontohan. Peluncuran satu inisiatif berskala kecil untuk menguji penggunaan AI dalam struktur organisasi tertentu dapat mengungkap tantangan yang mungkin timbul selama transformasi dengan cara yang aman dan terkendali. Menerapkan strategi ini dapat menghemat banyak kesulitan di masa mendatang.

Ketika mempertimbangkan penggunaan AI untuk tujuan manajemen risiko, perlu diingat bahwa saat ini algoritma AI ditugaskan untuk melakukan operasi tertentu dengan cakupan terbatas. Pengenalan visual atau ucapan dan analisis teks dapat menjadi contoh. Meskipun untuk tugas khusus tersebut AI memang mampu mengungguli manusia, AI yang serba guna dan berkemampuan tak terbatas sebagian besar masih dalam tahap konseptual. Karena alasan ini, penciptaan ekosistem hibrida berdasarkan kombinasi metode komputasi dan kecerdasan manusia kemungkinan akan menghasilkan hasil terbaik. Saat ini, pengambilan keputusan berbasis risiko membutuhkan dua elemen yang berbeda namun saling terkait erat: data objektif dan perasaan subjektif tentang apa yang akan dicapai atau hilang dari keputusan tersebut. Kalkulasi objektif dan persepsi subjektif sangat penting; keduanya tidak cukup dengan sendirinya. Sebaliknya, bagi mesin, pilihan manusia sebagian besar merupakan hasil

dari penggunaan keahlian dan keterampilan yang diperoleh secara otomatis dan berulang, alih-alih perenungan, evaluasi, dan penilaian yang matang. Dalam kondisi yang belum pernah terjadi sebelumnya dan sensitif terhadap waktu, individu sering kali perlu mengambil keputusan tanpa banyak pertimbangan. Dalam hal ini, 'kerja sama manusia-mesin' tampaknya menjadi cara terbaik. Melalui penggunaan AI, kendala yang ada dalam proses manajemen risiko dapat diatasi, dan keputusan yang didasarkan pada intuisi akan menjadi berbasis data dan sistematis. AI yang tersedia 24/7 dan 100 persen mutakhir dapat memungkinkan pemahaman risiko yang lebih jelas, memberikan informasi dan dukungan yang tepat waktu kepada staf penerbangan yang ditunjuk, dan dengan demikian dapat secara signifikan meningkatkan peluang sektor ini untuk mengambil langkah-langkah korektif dan preventif tepat waktu.

Dalam beberapa tahun mendatang, AI diharapkan dapat mendukung dan meningkatkan kinerja berbagai pelaku penerbangan dengan membantu di area-area yang kemampuannya jauh lebih besar daripada kemampuan operator manusia, seperti deteksi dan identifikasi risiko baru, manajemen armada dan staf, manajemen lalu lintas udara, dan lainnya. Meskipun kemampuannya masih terbatas, dan masih banyak tantangan yang bersifat regulasi atau etika, industri penerbangan harus didorong untuk meletakkan fondasi bagi penerapan teknologi ini sejak dini, dengan mengambil satu langkah kecil pada satu waktu.

8.5 RESPONS GLOBAL ATAS KERANGKA HUKUM

AI terus-menerus dibangun ke dalam mesin penerbangan komersial yang rumit, dan proses ini kemungkinan akan semakin cepat. Oleh karena itu, respons terkoordinasi secara global dengan tujuan menetapkan prinsip-prinsip yang mengatur pemanfaatannya telah menjadi hal yang mendesak. Agak merepotkan bahwa, meskipun penggunaan AI dalam bisnis penerbangan semakin meluas dan perhatian yang telah diperolehnya di antara para pembuat keputusan, respons internasional terhadap masalah ini masih kurang.

Maskapai penerbangan, ketika melakukan penerbangan internasional, seringkali tunduk pada berbagai undang-undang yang seringkali kontradiktif, yang mengatur berbagai bidang operasi mereka. Ini mencakup hal-hal seperti hak penumpang dan tanggung jawab maskapai penerbangan, serta standar keselamatan kesehatan dan langkah-langkah perlindungan yang harus diikuti ketika terbang ke negara yang ditentukan. Bagi bisnis yang bersifat global, peraturan yang rumit seperti itu sangat bermasalah. Demikian pula, dampak AI melampaui entitas tertentu atau bahkan melampaui batas sektor mana pun. Oleh karena itu, penggunaan solusi AI dapat menimbulkan dampak disruptif bagi mereka yang terlibat dalam aktivitas penerbangan, seperti maskapai penerbangan, produsen, bandara, penyedia layanan navigasi udara, atau penyedia layanan pemeliharaan. Tentu saja, hal ini juga akan berdampak signifikan pada lingkaran yang lebih luas. Oleh karena itu, terdapat kebutuhan yang kuat untuk mendorong kolaborasi global dan mengkonsolidasikan upaya-upaya industri yang terisolasi, regional, atau sektoral. Kompleksitas permasalahan ini membutuhkan keterlibatan beragam aktor dari berbagai latar belakang dan keahlian yang luas untuk berkolaborasi dalam kebijakan-

kebijakan yang relevan serta menangani masalah-masalah publik dan etika yang berkaitan dengan AI.

Sebaiknya, tindakan-tindakan tersebut dilakukan di bawah naungan Organisasi Penerbangan Sipil Internasional (ICAO), sebuah badan khusus Perserikatan Bangsa-Bangsa, yang dibentuk oleh Konvensi Chicago tahun 1944 dan dipercayakan dengan tugas-tugas mengembangkan 'prinsip dan teknik navigasi udara internasional' dan mempromosikan 'perencanaan dan pengembangan transportasi udara internasional'.

Mengingat beragamnya solusi AI dan potensi aplikasinya, berbagai komentar dapat diajukan mengenai isu-isu hukum paling mendesak yang perlu ditangani setelah implementasi AI di sektor penerbangan. Namun, dua masalah khusus, yang dipercepat oleh krisis COVID-19, membutuhkan perhatian segera dari komunitas penerbangan sipil internasional. Yaitu, keselamatan penerapan AI dan tata kelola data di sektor penerbangan yang dilengkapi dengan solusi AI.

Sejak awal transportasi udara komersial, keselamatan menjadi perhatian utama bagi komunitas penerbangan. Maka, tidak mengherankan jika ICAO, sejak awal berdirinya, telah berfokus pada keselamatan dan memainkan peran penting dalam menetapkan aturan yang berfokus pada pengoperasian pesawat yang aman. Industri ini telah berkembang selama bertahun-tahun, dan konstruksi pesawat telah meningkat seiring dengan banyaknya undang-undang udara keras dan lunak yang membahas keselamatan penerbangan, dan selanjutnya keamanan penerbangan. Akibatnya, saat ini, tidak ada moda transportasi setara yang dapat dikaitkan dengan tingkat keselamatan yang begitu tinggi dan dengan prosedur naik pesawat yang begitu ketat seperti pesawat komersial. Implementasi AI tidak diragukan lagi akan mengubah lanskap risiko keselamatan dan keamanan. Mengingat hal di atas, penggunaan AI, sebagai solusi manajemen risiko dan risiko keselamatan dan keamanan yang muncul pada saat yang sama, sangat menarik bagi ICAO.

Perlu dicatat bahwa isu AI telah diangkat di forum ICAO pada beberapa kesempatan. Misalnya, selama Konferensi Navigasi Udara Ketigabelas, Singapura menyampaikan makalah berjudul 'Potensi Kecerdasan Buatan dalam Manajemen Lalu Lintas Udara'. Makalah tersebut menyoroti kemampuan AI dalam domain pengambilan keputusan, dan potensinya untuk menjamin bahwa penerbangan tidak terhalang oleh keterbatasan manusia. Selain itu, selama Sesi Majelis ke-40, Makalah Kerja tentang AI disampaikan oleh Dewan Koordinasi Internasional Industri Dirgantara (ICCAIA) dan Organisasi Layanan Navigasi Udara Sipil (CANSO), yang menyerukan amandemen dan tindakan Standar dan Praktik yang Direkomendasikan (SARP), khususnya di bidang-bidang seperti sertifikasi, operasi, kualifikasi, dan berbagi data dan pelatihan. Namun demikian, meskipun ICAO terlibat dalam berbagai inisiatif di bidang AI, termasuk partisipasi dalam Kelompok Kerja EUROCAE 114 (WG-114) tentang Kecerdasan Buatan, tidak ada indikasi yang jelas tentang tindakan apa pun yang dicoba untuk meninjau atau meningkatkan SARP yang relevan. Selain tindakan yang ditujukan untuk mengeksplorasi dan mempromosikan potensi AI, ICAO harus berupaya memastikan adanya kerangka kerja yang memadai untuk penggunaan AI yang aman dalam penerbangan komersial.

Selain masalah keselamatan yang telah disebutkan di atas, isu lain yang menuntut tindakan global adalah tata kelola data di sektor penerbangan. Tak perlu dikatakan lagi bahwa AI membutuhkan kualitas yang memadai dan kuantitas informasi yang tepat untuk menjalankan fungsinya guna menghindari efek GIGO ('sampah masuk, sampah keluar'). Namun, tata kelola data lebih dari sekadar manajemen kualitas dan kuantitas data; tata kelola data juga tentang memastikan bahwa data yang diperoleh andal dan tersedia saat dibutuhkan, bahwa infrastruktur yang tepat disediakan dan aman (termasuk dalam hal risiko siber), bahwa kerangka kerja berbagi data telah tersedia, dan bahwa isu-isu seperti potensi pelanggaran data ditangani dan dikelola. Karena tantangan yang terkait dengan pemrosesan data bukanlah hal baru bagi penerbangan internasional, diharapkan tantangan tersebut juga akan ditangani dalam konteks AI.

Karena sektor penerbangan hanya dapat memperoleh manfaat dari pendekatan yang terpadu dan terkoordinasi, kerja sama antara ICAO, negara-negara, pemangku kepentingan penerbangan, dan perwakilan industri untuk menetapkan standar AI, terutama di area yang telah dibahas di atas (yaitu keselamatan dan tata kelola data), sangatlah penting.

8.6 KESIMPULAN

Menghadapi kesulitan keuangan yang parah akibat COVID-19, implementasi AI mungkin tidak tampak sebagai kebutuhan dasar bagi industri ini. Tak diragukan lagi, penerbangan komersial perlu belajar dari kesalahan yang telah terjadi, dan tak diragukan lagi ini adalah langkah yang tepat dan terbukti di sektor penerbangan. Namun, hanya mengandalkan strategi sukses di masa lalu mungkin terbukti tidak hanya tidak efektif saat ini, tetapi juga berbahaya bagi masa depan penerbangan komersial. Industri penerbangan perlu memperluas perspektifnya melampaui tantangan menjalankan bisnis di lingkungan pandemi saat ini, melampaui penyelesaian masalah yang ditimbulkan COVID-19, tidak hanya agar industri ini dapat berkembang, tetapi juga untuk membatasi risiko terulangnya krisis besar di masa mendatang. Seperti yang dikatakan Profesor Brian Havel, Direktur Institute of Air and Space Law di McGill University: '[p]enerbangan internasional bukanlah bisnis, di mana masalah akan hilang jika kita membiarkannya begitu saja, di mana masalah dapat dibiarkan menunggu selamanya'. COVID-19 telah menempatkan industri transportasi udara di bawah tekanan dan menyerukan fleksibilitas dan ketangkasan yang lebih besar untuk menghadapi tantangan ketidakpastian yang telah berlangsung lama. Oleh karena itu, komunitas penerbangan berupaya untuk menerapkan berbagai perbaikan menggunakan teknologi yang tersedia. AI, sebagai salah satu kemajuan paling menjanjikan dalam hal identifikasi risiko, membutuhkan perhatian khusus. AI adalah teknologi yang mengganggu dan merupakan risiko tersendiri. Setelah percepatan implementasinya di sektor penerbangan, masalah seperti keselamatan dan tata kelola data perlu ditangani tanpa penundaan. Saat ini, menangani masalah AI dalam industri penerbangan merupakan hal yang mendesak baik dari perspektif regulasi maupun bisnis. Kerangka kerja yang tepat diperlukan untuk membangun kepercayaan dan mengatasi persepsi negatif terhadap AI, mempertahankan dan meningkatkan tingkat keselamatan dalam penerbangan sipil, mengurangi liabilitas dan tantangan terkait asuransi, serta memberikan

tingkat kepastian yang dapat diterima bagi bisnis terkait penggunaan AI. Oleh karena itu, respons global yang terkoordinasi yang melibatkan beragam pakar AI dan pelaku industri diperlukan untuk memastikan tercapainya tujuan-tujuan ini. Upaya Eropa di bidang AI, termasuk inisiatif EASA dan EUROCONTROL, dapat menjadi contoh inspiratif dan landasan yang kokoh bagi diskusi internasional yang akan diselenggarakan di bawah kepemimpinan ICAO.

BAB 9

AI OTONOM, PELABUHAN PINTAR, DAN MANAJEMEN RANTAI PASOK

9.1 AI, PELABUHAN PINTAR, DAN KAPAL CERDAS

Tugas menangani kargo yang terus meningkat jumlahnya dengan cara yang aman, efisien, dan ramah lingkungan merupakan salah satu tantangan terbesar yang dihadapi pelabuhan dan industri maritim secara luas. Literatur mendefinisikan "pelabuhan" sebagai jaringan logistik tempat arus kargo, uang, dan informasi saling terkait dengan berbagai aspek politik, ekonomi, sosial, teknologi, dan lingkungan. Jaringan logistik semacam itu menggunakan AI untuk mengotomatiskan area operasi, meningkatkan komunikasi, dan profitabilitas. Aktivitas pelabuhan memang semakin berteknologi. Saat ini, Internet of Things (IoT) secara khusus dianggap sebagai tambahan teknologi yang penting.

Di sisi lain, "pelabuhan pintar" dapat didefinisikan sebagai pelabuhan yang sepenuhnya otomatis tempat semua perangkat terhubung melalui IoT. Yang dkk. Bahasa Indonesia: menawarkan contoh proyek di pelabuhan Rotterdam, Hamburg, Le Havre, Shanghai, Vigo, dan Singapura yang berkaitan dengan pelabuhan pintar dan IoT. Ini termasuk teknologi penginderaan, derek dermaga otomatis, kendaraan berpemandu otomatis untuk penanganan kontainer dan derek halaman. Jaringan sensor dan aktuator pintar, perangkat nirkabel dan pusat data membentuk infrastruktur utama pelabuhan pintar, yang memungkinkan otoritas pelabuhan untuk menyediakan layanan penting dengan cara yang lebih cepat dan lebih efisien. Selain itu, menurut beberapa perkiraan, terminal kontainer otomatis dapat menghemat setidaknya 25% lebih banyak energi dan mengurangi 15% lebih banyak emisi karbon daripada terminal tradisional. Dengan demikian, sistem pelabuhan mengalami transformasi yang mendalam, menerapkan teknologi yang terkait dengan sensor dan aktuator pintar terdistribusi, komunikasi data dan konektivitas internet untuk operasi jarak jauh dan otomatis, optimalisasi kontrol berdasarkan AI, serta analisis data besar. Pelabuhan pintar tersebut menggunakan sistem informasi terkait AI yang "mengelola, memantau, dan menyimpan sejumlah besar data (misalnya lalu lintas maritim dan data logistik) dan menyediakan Bahasa Indonesia: layanan terkomputerisasi dan tanpa kertas berskala besar di pelabuhan pintar". Keragaman data dan informasi yang dikumpulkan kemudian memungkinkan "aplikasi pelabuhan pintar untuk beradaptasi dengan persyaratan dinamis dari sistem kompleks yang menangani beragam aspek dan terdiri dari berbagai teknologi (misalnya komunikasi nirkabel dan sistem tertanam untuk operasi penginderaan)".

Kemunculan COVID-19 yang tak terduga telah memberikan tantangan dan peluang yang mendesak pada sektor maritim dan penggunaan AI. Karantina terkait COVID-19, penutupan pelabuhan, pemeriksaan keamanan di pelabuhan yang mengakibatkan waktu tunggu tambahan untuk operasi berlabuh, operasi transshipment pelabuhan laut pedalaman, manajemen rantai pasokan, dan aspek-aspek lainnya berdampak pada jumlah keseluruhan

barang dan pengiriman laut. Penggunaan pelabuhan pintar dan kapal pintar yang lebih besar dapat mengurangi dampak buruk tersebut, dan digitalisasi terkait AI dapat menjadi warisan pandemi.

Secara analitis, seluruh ekosistem pelabuhan pintar terdiri dari lima area aplikasi utama: (a) manajemen kapal pintar; (b) manajemen kontainer pintar; (c) manajemen pelabuhan pintar; (d) manajemen energi pintar; dan (e) manajemen sumber daya cerdas. Pelabuhan cerdas yang diatur AI memiliki potensi untuk mengurangi kesalahan manusia, memungkinkan komunikasi aktif dengan lingkungan sosial, membuat operasi lebih cepat, lebih aman, dan lebih efisien, serta mengurangi emisi karbon untuk keberlanjutan lingkungan dan energi. Selain itu, pelabuhan cerdas menawarkan solusi adaptif dalam menanggapi dunia yang dicirikan oleh ketidakpastian. Sistem pendukung pengambilan keputusan AI tersebut berdasarkan model perilaku prediktif menghemat biaya transaksi, menerapkan manajemen risiko yang efisien, meningkatkan keselamatan pekerja dan memungkinkan alokasi sumber daya yang sangat efisien, daya saing, dan pembangunan berkelanjutan. Akhirnya, AI dapat berkontribusi pada "kapal pintar", yaitu sistem dengan navigasi, pelayaran, kontrol, dan panduan yang sepenuhnya otonom. Ini kadang-kadang disebut sebagai "kapal cerdas".

9.2 KECERDASAN BUATAN PELABUHAN PINTAR DAN CRM DI INDONESIA

Penerapan kecerdasan buatan (*Artificial Intelligence/AI*) dan sistem otonom di sektor pelabuhan dan rantai pasok di Indonesia menunjukkan potensi yang sangat besar dalam meningkatkan efisiensi dan keamanan operasi logistik. Teknologi AI memungkinkan otomatisasi proses bongkar-muat, pengaturan arus kontainer, serta pemantauan inventaris secara real-time, sehingga pelabuhan pintar dapat beroperasi dengan waktu tunggu kapal yang lebih singkat dan akurasi pergerakan barang yang lebih tinggi. Contoh nyata di Indonesia dapat dilihat dari Pelabuhan Tanjung Priok di Jakarta, pelabuhan terbesar dan tersibuk di Indonesia, yang telah mulai mengimplementasikan sistem digital berbasis AI. Sistem ini digunakan untuk memprediksi pergerakan kontainer, mengoptimalkan jalur kendaraan otomatis di area pelabuhan, dan meminimalkan kesalahan manusia. Dampaknya terlihat pada meningkatnya kecepatan bongkar-muat kapal dan penurunan waktu tunggu, sekaligus mengurangi risiko kerusakan barang akibat kesalahan operasional. Di sisi manajemen rantai pasok, AI digunakan untuk memprediksi permintaan barang, mengoptimalkan rute transportasi, serta memantau kondisi inventaris di berbagai titik distribusi. Salah satu implementasi nyata dapat dilihat pada perusahaan logistik nasional seperti JNE dan Pelindo, yang mulai mengintegrasikan sistem AI untuk memantau arus pengiriman, memprediksi kepadatan jalur transportasi, dan mengoptimalkan rute distribusi. Dengan kemampuan analisis data yang tinggi, AI mampu mengidentifikasi potensi gangguan atau penundaan dalam rantai pasok dan memberikan rekomendasi mitigasi secara cepat. Hal ini mengurangi risiko kekurangan atau kelebihan stok di titik distribusi dan memperbaiki keandalan sistem logistik nasional.

Meskipun demikian, penerapan AI otonom dalam pelabuhan pintar dan rantai pasok juga menimbulkan sejumlah tantangan dan risiko. Dari sisi teknologi, integrasi antara sistem AI

dengan infrastruktur lama di pelabuhan masih menjadi kendala, terutama terkait kompatibilitas data dan interoperabilitas perangkat. Risiko keamanan siber menjadi isu penting karena seluruh proses operasional dan data logistik tersambung secara digital, sehingga potensi gangguan atau peretasan sistem dapat berdampak besar terhadap operasi pelabuhan dan rantai pasok. Selain itu, dari sisi regulasi, masih terdapat kekosongan hukum yang jelas mengenai tanggung jawab ketika terjadi kesalahan atau kerusakan akibat sistem AI. Tantangan lainnya adalah adaptasi sumber daya manusia, di mana operator dan pekerja pelabuhan perlu dilatih untuk bekerja sama dengan sistem otonom agar teknologi AI dapat dimanfaatkan secara optimal tanpa menimbulkan gangguan operasional.

Dengan demikian, pemanfaatan AI otonom di pelabuhan pintar dan manajemen rantai pasok di Indonesia menawarkan peluang besar untuk meningkatkan produktivitas, keselamatan, dan keandalan logistik nasional. Keberhasilan implementasinya terlihat dari percepatan operasional di Pelabuhan Tanjung Priok dan peningkatan efisiensi distribusi di perusahaan logistik nasional, namun tetap membutuhkan kesiapan teknologi, dukungan regulasi yang jelas, serta kemampuan sumber daya manusia dalam mengelola sistem baru ini. Pendekatan yang seimbang antara inovasi teknologi dan manajemen risiko diperlukan agar manfaat AI dapat direalisasikan secara maksimal tanpa menimbulkan konsekuensi negatif bagi operasi logistik dan ekonomi nasional.

9.2 INTEGRASI DAN KOLABORASI

Pelabuhan laut dengan ekosistem AI dibangun di sekitar "sistem siber-fisik", yaitu infrastruktur fisik dengan fasilitas siber yang memengaruhi efisiensi industri maritim di berbagai tingkatan, termasuk kapal, pelabuhan, dan rantai pasokan terkait. AI di bidang manajemen rantai pasokan digunakan untuk menemukan pola dalam rantai logistik dan menawarkan prediksi waktu terperinci kapan kapal, truk, dan kontainer akan mencapai terminal, sehingga memungkinkan perencanaan yang lebih baik. Selain itu, AI dapat digunakan untuk memprediksi kebutuhan peralatan di masa mendatang, pemanfaatan lapangan penumpukan jangka panjang, kerusakan kontainer, jumlah kunjungan gerbang, dan banyak aspek lainnya.

Pelabuhan laut di Rotterdam, Hamburg, Le Havre, Shanghai, Vigo, dan Singapura, misalnya, telah membangun apa yang disebut "integrasi rantai pasokan". Istilah ini merujuk pada "berbagi data, integritas, dan kemitraan transparan dengan teknologi informasi, integrasi sistem, dan kolaborasi sebagai ideologi pemandu, dan berpusat pada perusahaan pelabuhan sebagai intinya". Proses terkait lainnya adalah apa yang disebut "struktur rantai pasokan perusahaan", yang dibentuk oleh "perusahaan hulu dan hilir terkait untuk menyediakan aliran informasi, logistik, dan modal yang lancar di seluruh rantai pasokan dari sumber tengah hingga pengiriman uang". Misalnya, manajemen pasokan kapal pintar mengawasi kapal, termasuk pilihan rute dan pelabuhan mereka, berdasarkan lokasi dan lalu lintas di dalam pelabuhan. Hal ini untuk meningkatkan ketepatan waktu kedatangan mereka di pelabuhan. Di sisi lain, manajemen kontainer pintar membantu akuisisi, pelacakan, transportasi, penyimpanan, dan reposisi kontainer, serta transshipment (pemindahan dari satu kapal ke kapal lain). Selain itu,

manajemen energi pintar mengurangi konsumsi energi. Sistem manajemen yang dioperasikan AI semacam itu dapat menyediakan jadwal yang efektif, alokasi sumber daya, dan optimalisasi dalam hal waktu dan biaya.

Ekosistem yang didukung AI meningkatkan kolaborasi antara berbagai pihak, seperti otoritas pelabuhan, pemilik kargo, dan penyedia logistik pihak ketiga, karena menyelaraskan peta jalan digital masing-masing. Hal ini pada gilirannya memungkinkan peluang yang saling menguntungkan untuk meningkatkan efisiensi dan mengurangi pemborosan (dengan berbagi data dan menyediakan antarmuka standar untuk wawasan, prediksi, dan kendala yang didukung AI). Selain itu, AI memberikan perkiraan yang jauh lebih akurat dan andal untuk kedatangan dan keberangkatan kapal hingga 80%, yang sangat meningkatkan operasi manajemen rantai pasokan. AI juga menganalisis dan memantau kargo dan kendaraan laut secara real-time. Ini mengotomatiskan analisis manual yang menantang, mempercepat perencanaan logistik kargo, dan meningkatkan deteksi situasi yang berpotensi luar biasa.

9.3 POTENSI RISIKO DAN MASALAH REGULASI

Konektivitas siber AI berskala besar dapat menimbulkan banyak tantangan keamanan. Misalnya, laporan tentang ancaman siber di sektor maritim yang disusun oleh Badan Keamanan Informasi dan Jaringan Eropa (ENISA) menunjukkan bahwa ancaman siber merupakan masalah serius. Salah satu isu yang paling memicu adalah cara terbaik untuk menerapkan keamanan siber dan cara menjaga keamanan pelabuhan digital dan otomatis. Pelabuhan dan kapal pintar yang dioperasikan AI merupakan subjek potensial serangan yang dapat mengakibatkan risiko keselamatan dan keamanan. Misalnya, menyusup ke komputer pelabuhan, mengirimkan sinyal GPS palsu untuk mengubah rute kapal, mengubah sinyal sistem identifikasi otomatis kapal untuk melaporkan lokasi mereka secara keliru, mengakses tampilan peta elektronik dan perangkat lunak sistem informasi untuk mengubah peta, semuanya merupakan masalah nyata dan dapat menimbulkan konsekuensi yang buruk. Risiko-risiko ini dapat berdampak pada potensi tanggung jawab atas kerugian yang ditimbulkan. Peraturan khusus industri yang terkait dengan subjek ini mencakup ketentuan yang berkaitan dengan standar keselamatan dan keamanan, sebagaimana tercantum dalam Kode Keamanan Kapal dan Fasilitas Pelabuhan Internasional (ISPS) dan dalam Peraturan Uni Eropa (EC) No 725/2004 tentang peningkatan keamanan kapal dan fasilitas pelabuhan. Namun, peraturan ini tidak menganggap serangan siber sebagai kemungkinan ancaman atau tindakan melanggar hukum dan tidak memadai untuk mengatasi potensi risiko terkait AI bagi keamanan publik.

Dalam hal potensi tanggung jawab atas kerusakan yang diakibatkan oleh penggunaan teknologi digital yang baru muncul, hukum Negara Anggota UE tidak memuat aturan tanggung jawab yang secara khusus berlaku untuk kerusakan tersebut. Masalah tanggung jawab di UE sebagian besar tidak terkoordinasi, dengan pengecualian hukum tanggung jawab produk berdasarkan Direktif 85/374/EC, beberapa aspek tanggung jawab atas pelanggaran hukum perlindungan data, dan tanggung jawab atas pelanggaran hukum persaingan. Untuk sebagian besar ekosistem teknologi modern, seperti pelabuhan pintar, kapal pintar, dan manajemen rantai pasokan AI terkait, tidak ada rezim tanggung jawab khusus. Beberapa orang

menyarankan bahwa aturan hukum gugatan domestik umum berlaku. Bergantung pada yurisdiksinya, ini termasuk aturan (atau aturan) yang memperkenalkan tanggung jawab berbasis kesalahan dengan cakupan aplikasi yang relatif luas, disertai dengan beberapa aturan yang lebih spesifik. Aturan-aturan khusus ini memodifikasi premis pertanggungjawaban berbasis kesalahan (terutama distribusi beban pembuktian kesalahan) atau menetapkan pertanggungjawaban yang independen dari kesalahan (biasanya disebut pertanggungjawaban ketat atau pertanggungjawaban berbasis risiko), yang juga mengambil banyak bentuk yang bervariasi sehubungan dengan ruang lingkup aturan, kondisi pertanggungjawaban, dan beban pembuktian. Selain itu, sebagian besar rezim pertanggungjawaban mengandung gagasan pertanggungjawaban untuk orang lain (yaitu pertanggungjawaban perwakilan, di mana seseorang dianggap bertanggung jawab atas tindakan orang lain).

Yurisdiksi yang telah mengizinkan penggunaan eksperimental atau reguler kendaraan yang sangat atau sepenuhnya otomatis biasanya menyediakan pertanggungan atas kerusakan yang disebabkan, baik melalui asuransi atau dengan mengacu pada aturan hukum umum. Misalnya, dalam sistem hukum Inggris, Undang-Undang Kendaraan Otomatis dan Listrik (AEVA 2018) (c 18) mengatur dalam s. 2 bahwa "perusahaan asuransi bertanggung jawab" atas kerugian yang diderita oleh tertanggung atau orang lain dalam kecelakaan yang disebabkan oleh kendaraan otomatis. Selain itu, Undang-Undang ini juga mewajibkan pendaftaran kendaraan otomatis yang dirancang atau diadaptasi agar mampu mengemudi sendiri oleh Departemen Perhubungan. Namun, AEVA 2018 masih menyisakan beberapa permasalahan substantif, seperti cara perusahaan asuransi dapat mengajukan klaim kembali dari produsen dan apakah terdapat pembelaan yang tersedia, serta terkait definisi. Lebih lanjut, beberapa pihak berpendapat bahwa terlepas dari undang-undang ini, dampak buruk dari pengoperasian teknologi digital yang sedang berkembang di Inggris dapat dikompensasi berdasarkan undang-undang ("tradisional") yang berlaku tentang ganti rugi dalam kontrak dan perbuatan melawan hukum.

9.4 MENUJU INTERVENSI REGULASI YANG OPTIMAL

Dalam kasus kerusakan yang disebabkan oleh AI dalam industri maritim, serupa dengan sektor lainnya, mungkin mustahil untuk mengidentifikasi pihak yang bertanggung jawab dalam memberikan kompensasi. Hal ini disebabkan oleh sifat dunia maya, serta kompleksitas ekosistem ini. Pentingnya keamanan siber di sektor maritim disebutkan dalam laporan ENISA. ENISA merekomendasikan agar pemerintah mengambil langkah-langkah yang tepat untuk menambahkan pertimbangan dan ketentuan terkait keamanan siber dalam kerangka regulasi maritim nasional. Secara mendalam, ENISA mendukung penerapan standar dan praktik keamanan yang ditingkatkan untuk keamanan siber di sektor maritim. Pendekatan semacam itu harus melibatkan kerja sama internasional dan keterlibatan yang kuat dengan sektor swasta, tetapi tidak boleh menempatkan pemerintah dalam posisi untuk menentukan desain dan pengembangan teknologi yang terlibat di masa depan.

Berdasarkan hukum Uni Eropa, Direktif 85/374/EEC hanya mencakup kerusakan yang disebabkan oleh cacat produksi AI dan dengan syarat bahwa orang yang terluka dapat

membuktikan hubungan sebab akibat antara kerusakan aktual dan cacat pada produk (dalam hal ini, AI). Selain itu, Direktif tersebut menangani aspek paling dasar dari sebab akibat dan menyerahkan sebagian besar masalah kepada hukum nasional. Hal ini berpotensi menyebabkan hasil yang berbeda terkait dengan peraturan pertanggungjawaban terkait teknologi AI. Direktif tersebut memang mencakup beberapa pembelaan, yang dapat ditemukan dalam pasal. 7. Misalnya, produsen tidak bertanggung jawab jika dapat dibuktikan bahwa "dengan mempertimbangkan keadaan, kemungkinan besar cacat yang menyebabkan kerusakan tersebut tidak ada pada saat produk tersebut diedarkan olehnya atau bahwa cacat tersebut muncul setelahnya". Lebih lanjut, produsen tidak akan bertanggung jawab jika "kondisi pengetahuan ilmiah dan teknis pada saat ia mengedarkan produk tidak memungkinkan ditemukannya cacat tersebut". Ketentuan-ketentuan ini memang bermakna, tetapi tidak sepenuhnya mencerminkan permasalahan teknologi yang disebutkan dalam bab ini. Secara khusus, rezim tanggung jawab produk saat ini beroperasi dengan asumsi bahwa produk tidak terus berubah dengan cara yang tidak dapat diprediksi dan tidak dapat diramalkan setelah meninggalkan jalur produksi. Seperti yang disarankan Infantino dan Zervogianni, salah satu faktor terpenting untuk menetapkan tanggung jawab dalam banyak sistem hukum Eropa adalah dapat diperkirakan, yang didasarkan pada gagasan bahwa "terdakwa harus bertanggung jawab hanya atas kerusakan yang orang yang wajar dengan kehati-hatian biasa yang ditempatkan pada posisinya akan dapat diperkirakan sebagai akibat yang mungkin terjadi dari tindakannya". Ini tidak selalu terjadi dengan teknologi AI, karena AI dapat menghasilkan solusi yang mungkin tidak dipertimbangkan oleh manusia. Menurut Martin-Casals, "dapat diperkirakan menghadirkan tantangan yang menjengkelkan bagi sistem hukum mana pun yang ingin memecahkan masalah pemberian ganti rugi kepada korban kerugian yang disebabkan oleh AI". Dia kemudian menyarankan bahwa mungkin pendekatan yang lebih baik adalah dengan mengandalkan teori yang didasarkan pada "cakupan risiko yang diciptakan oleh aktivitas tergugat". Teori "kerugian dalam risiko" ini sebenarnya tidak asing bagi praktik Eropa, karena dapat ditemukan sebagai dasar dari beberapa pengujian yang sudah diterapkan di Eropa. Selain itu, Kötz dan Wagner menyarankan bahwa tujuan perlindungan dari aturan semacam itu (yang mengacu pada kerugian yang merupakan efek berkelanjutan dari risiko yang membuat pelaku kesalahan bertanggung jawab), bersama dengan risiko umum kehidupan (yang dapat dikaitkan dengan korban) memungkinkan solusi yang lebih baik tanpa memerlukan spekulasi yang terlibat dalam uji keterlihatan. AI adalah sumber risiko yang melebihi risiko lain yang bersumber dari perilaku manusia. Jadi, bahkan dalam kasus ketika kerugian disebabkan oleh perilaku AI yang tidak dapat diramalkan, pemrogramnya juga harus bertanggung jawab.

Literatur hukum dan ekonomi menawarkan beberapa tanggapan menarik terhadap potensi intervensi regulasi di bidang AI, yang juga dapat diterapkan pada industri maritim. Misalnya, Shavell berpendapat bahwa jika terdapat pihak (prinsipal) yang memiliki kendali atas perilaku pihak lain (agen/AI), maka prinsipal dapat dianggap bertanggung jawab secara vikaris atas kerugian yang disebabkan oleh agen tersebut. Oleh karena itu, tanggung jawab vikaris dan hubungan prinsipal-agen yang spesifik dapat diperkenalkan. Prinsipal dapat dianggap

bertanggung jawab secara vikaris atas kerusakan dan/atau kerugian yang disebabkan oleh agen (AI otonom). Perluasan tanggung jawab ini secara tidak langsung dapat mengarah pada pengurangan risiko dan standar industri yang lebih tinggi. Secara khusus, hal ini dapat memberlakukan beberapa standar berikut: a) pertanggungan asuransi wajib (misalnya oleh prinsipal); b) peraturan mengenai standar keselamatan; c) subsidi untuk tindakan pencegahan dan denda karena menyebabkan kerugian; dan d) pembentukan dana asuransi yang didanai publik-swasta di seluruh dunia.

Selain itu, beberapa pihak mengusulkan penerapan tanggung jawab mutlak bagi produsen AI, yang dapat dilengkapi dengan persyaratan bahwa pelanggaran standar keselamatan wajib yang tidak dapat dimaafkan merupakan *kelalaian semata*. Namun, kepatuhan terhadap standar ini tidak seharusnya menghapuskan tanggung jawab atas perbuatan melawan hukum.

9.5 KESIMPULAN

Tugas menangani kargo dalam jumlah yang terus meningkat dengan cara yang aman, efisien, dan ramah lingkungan merupakan salah satu tantangan terbesar yang dihadapi pelabuhan dan industri maritim secara luas. Pelabuhan pintar yang dikelola AI berpotensi mengurangi kesalahan manusia, memungkinkan komunikasi aktif dengan lingkungan sosial, mempercepat, mengamankan, dan mengoptimalkan operasi, serta mengurangi emisi karbon demi keberlanjutan lingkungan dan energi. Pelabuhan pintar tersebut berupaya menyediakan manajemen rantai pasokan yang lancar, mengintegrasikan sisi penawaran dan permintaan untuk mengoptimalkan alokasi sumber daya, layanan, dan pengawasan yang relevan, serta bongkar muat secara otonom.

Namun, penggunaan AI yang berlebihan dapat menimbulkan berbagai potensi risiko baru terkait keamanan maritim dan rantai pasokan. Ketika kapal dan pelabuhan menjadi semakin "cerdas" dan membutuhkan semakin sedikit orang untuk campur tangan dalam aktivitas mereka, pertanyaan bagi para pembuat undang-undang adalah bagaimana mencegah, menghalangi, dan memitigasi potensi risiko publik yang ditimbulkan oleh AI. Misalnya, salah satu isu yang paling memicu adalah bagaimana cara terbaik menerapkan keamanan siber dan bagaimana menjaga keamanan pelabuhan digital dan otomatis. Pelabuhan dan kapal pintar yang dioperasikan AI berpotensi menjadi sasaran serangan yang dapat mengakibatkan risiko keselamatan dan keamanan. Akibatnya, data yang salah dapat mengakibatkan kerugian yang tidak terduga, bahaya keselamatan publik, dan masalah pertanggungjawaban yang memerlukan penanganan regulasi yang substantif.

Secara umum, mekanisme hukum saat ini secara memadai menangani tanggung jawab atas AI non-otonom yang diawasi manusia. Namun, jika AI berevolusi dan menjadi tanpa pengawasan, maka intervensi regulasi diperlukan. Baik ISPS maupun Panduan Keamanan Maritim dari Organisasi Maritim Internasional (IMO) harus diamandemen untuk memasukkan risiko spesifik terkait AI. Demikian pula, Peraturan (EC) No. 725/2004 tentang peningkatan keamanan kapal dan fasilitas pelabuhan saat ini dapat diamandemen untuk memasukkan risiko terkait AI. Di antara amandemen yang diusulkan yang disebutkan dalam bab ini adalah

dimasukkannya cakupan asuransi wajib, tetapi juga penerapan standar keselamatan, dan pembentukan dana asuransi yang didanai publik dan swasta.

BAB 10

AI DALAM KEBIJAKAN IKLIM DAN ENERGI UE–JEPANG

10.1 PENDAHULUAN

Program ilmiah dan teknologi memiliki dampak jangka panjang yang semakin meningkat. Proses teknologi yang revolusioner dan tampak visioner seperti yang terlibat dalam AI ... sedang menjadi kenyataan. Kita harus bersiap sekarang untuk realitas masa depan. Meningkatnya durasi proyek teknologi individual, terutama yang besar seperti penelitian bentuk-bentuk energi baru, membuat pertimbangan perspektif jangka panjang dari keputusan penelitian saat ini dan dampaknya terhadap generasi mendatang menjadi semakin penting.

Meskipun berasal dari tahun 1970-an, kutipan di atas tidak kehilangan signifikansinya di abad ke-21. Namun, karena kemajuan luar biasa yang telah dicapai selama bertahun-tahun menuju kemajuan AI, Internet of Things (IoT), robotika, dan pembelajaran mendalam, inovasi-inovasi yang sebelumnya inovatif telah menjadi barang sehari-hari. Revolusi AI memengaruhi hampir setiap aspek kehidupan kita, karena AI menemukan kegunaannya di berbagai cabang dan industri.

Proses ini tidak mengabaikan sektor energi. Sebaliknya, AI telah memberikan dampak yang berkelanjutan pada pasar energi, secara bertahap memengaruhi iklim dan kebijakan energi. Misalnya, AI dapat mendukung modernisasi jaringan listrik, meningkatkan stabilitasnya, dan mencegah pemadaman listrik. Hal ini dapat dicapai dengan mengendalikan pasokan dan permintaan di tingkat lokal dan nasional, membantu konsumen energi menghemat energi dan uang dengan memungkinkan pemantauan penggunaan mereka sendiri secara real-time dan menyesuaikan tarif yang sesuai, serta memfasilitasi interaksi dengan jaringan. Dengan demikian, AI dapat meningkatkan efisiensi energi dan mengatasi kemiskinan energi, memperkuat produksi energi dari sumber terbarukan (juga oleh individu), dan memfasilitasi pengurangan emisi.

Dalam konteks ini, bab ini mengkaji AI melalui prisma kebijakan energi dan aksi iklim Uni Eropa dan Jepang. Kedua contoh ini dipilih untuk dianalisis karena Uni Eropa dan Jepang memiliki pengalaman luas dalam penelitian dan pengembangan. Pengalaman ini, jika dipadukan dengan rencana energi dan ambisi iklim mereka, termasuk mencapai emisi gas rumah kaca nol bersih pada tahun 2050, dapat digunakan untuk menyoroti kebutuhan transisi energi bagi negara-negara ekonomi dunia yang penting. Dengan latar belakang tersebut, studi ini mencakup pendekatan publik terhadap hal-hal tersebut beserta aspek regulasi terkait. Bahasa Indonesia: Ini menyandingkan dokumen-dokumen utama, baik yang lalu maupun yang sekarang, yang relevan dengan analisis AI dalam bidang iklim dan energi dari negara-negara ekonomi terkemuka ini. Dalam konteks ini, UE berencana untuk membuka kerangka keuangan multitalitahunan barunya (2021–2027) untuk investasi dalam pembelajaran mesin tanpa pengawasan, energi, dan efisiensi data, yang menawarkan dukungan untuk teknologi dan infrastruktur yang lebih hemat energi yang 'membuat rantai nilai AI lebih hijau'. Di Jepang, data berkualitas tinggi telah digunakan selama bertahun-tahun untuk meningkatkan

produktivitas di lokasi produksi monozukuri. Pada abad ke-21, AI, IoT, big data, dan robotika disebut sebagai 'kunci terpenting untuk memimpin revolusi masa depan dalam produktivitas' di Jepang, serta 'revolusi industri keempat', yang Jepang tempatkan di jantung strategi revitalisasinya. 'Revolusi industri baru' dari ekonomi berbasis data (membentuk kembali sektor bisnis dengan energi sebagai elemennya) juga menjadi pusat perhatian Eropa.

Selain itu, bab ini Memperkenalkan pembahasan tentang penerapan AI di sektor energi listrik. Hal ini menyangkut kemungkinan penguatan regulasi energi di bawah model day-watchman dengan bantuan AI (menunjukkan seberapa besar AI dapat membantu regulator energi dalam tugas sehari-hari mereka). Hasilnya, gambaran yang lebih luas tergambar, yaitu menunjukkan bagaimana AI dapat diregulasi dalam bidang kebijakan iklim dan energi.

10.2 INSENTIF DAN KESEPAKATAN GLOBAL

Isu perlindungan lingkungan dan penanggulangan perubahan iklim sangat relevan bagi Uni Eropa. Mengingat dampak sektor energi terhadap lingkungan dan iklim (energi sebagai 'faktor kunci dalam pencapaian pembangunan berkelanjutan'), dapat ditemukan hubungan erat antara kedua bidang ini. Pendekatan semacam itu dapat disebut pro-lingkungan, yaitu dengan menjamin bahwa regulasi yang diterapkan oleh lembaga-lembaga Uni Eropa di sektor energi bertujuan untuk memastikan pelestarian lingkungan beserta sumber daya alamnya secara berkelanjutan. Kecenderungan ini telah menghasilkan pembentukan pendekatan bersama dalam kerangka kebijakan iklim-energi Uni Eropa.

Pertama, kebutuhan akan legislasi tentang isu-isu iklim secara resmi diakui oleh komunitas Eropa pada akhir 1980-an. Menanggapi aktivitas dan perkembangan internasional mengenai perubahan iklim, Komisi mengusulkan cara-cara tertentu untuk meluncurkan pendekatan Eropa untuk mengatasi efek rumah kaca. Kebijakan yang diusulkan meliputi: penelitian, tindakan pencegahan, adaptasi terencana, dan kerja sama dengan negara-negara berkembang. Selanjutnya, berbagai tindakan di bidang ini dilakukan pada 1990-an, yang merupakan langkah pertama menuju kebijakan iklim Eropa. Pada 1993, legislasi Eropa awal tentang perubahan iklim disahkan, dan pemantauan CO₂ dan emisi gas rumah kaca lainnya dimulai. Hal ini selanjutnya didorong oleh rezim internasional di mana UE telah berpartisipasi secara aktif – Konvensi Kerangka Kerja Perserikatan Bangsa-Bangsa tentang Perubahan Iklim diikuti oleh Protokol Kyoto untuk Konvensi ini.

Yang terakhir, meskipun ditolak oleh Amerika Serikat pada 2001, memotivasi UE untuk mengambil peran utama dalam bidang kebijakan iklim di tingkat internasional. Selain itu, berlakunya Protokol Kyoto Protokol pada tahun 2005 mendorong agenda iklim Eropa, yang mengarah pada pembentukan asumsi dasar aksi iklim Uni Eropa hanya dua tahun kemudian. Protokol ini memperkenalkan tiga pilar: pengurangan emisi gas rumah kaca, promosi sumber energi terbarukan, dan peningkatan efisiensi energi, semuanya dengan target persentase pan-Eropa: umumnya, '3 × 20%' untuk tahun 2020. Antara tahun 2008 dan 2009, proposal ini melewati proses legislatif dan disahkan, yang kemudian dikenal luas sebagai Paket Iklim dan Energi. Target yang ditetapkan untuk tahun 2020 direvisi pada tahun 2014 dengan meningkatkan pengurangan emisi menjadi setidaknya 40% dan pertumbuhan energi

terbarukan serta peningkatan efisiensi energi menjadi setidaknya 27%. Dua tujuan terakhir kembali ditingkatkan pada tahun 2018. Hal ini dicapai dengan mengadopsi arahan baru mengenai energi terbarukan dan efisiensi energi, yang menetapkan setidaknya 32% untuk pangsa sumber energi terbarukan dari konsumsi akhir bruto Uni Eropa dan mengurangi konsumsi energi Uni Eropa setidaknya sebesar 32,5% melalui peningkatan efisiensi energi pada tahun 2030.

Kemudian, pada bulan Desember 2015, perjanjian multilateral pertama tentang perubahan iklim yang mencakup hampir semua emisi global – Perjanjian Paris – ditandatangani. Di bawah rezim internasional ini, kontribusi yang ditentukan secara nasional (NDC) Uni Eropa ditetapkan di bawah kerangka iklim dan energi 2030 yang lebih luas (yang disebutkan sebelumnya 40%, 32%, 32,5%) untuk emisi, energi terbarukan, dan efisiensi energi. Pada akhir tahun 2019, didorong oleh aksi mogok iklim dan seruan untuk mendeklarasikan keadaan darurat iklim, Komisi Eropa mengumumkan peta jalan awal berisi kebijakan dan langkah-langkah utama yang diperlukan untuk mencapai Kesepakatan Hijau Eropa. Kesepakatan yang ditetapkan sebagai strategi pertumbuhan ini membawa tujuan-tujuan ambisius Uni Eropa, yaitu 'masyarakat yang adil dan sejahtera, dengan ekonomi yang modern, hemat sumber daya, dan kompetitif, yang tidak memiliki emisi gas rumah kaca bersih pada tahun 2050 dan yang pertumbuhan ekonominya dipisahkan dari penggunaan sumber daya'. Lebih lanjut, pada bulan Maret 2020, sebuah proposal regulasi tentang Hukum Iklim Eropa, yang bertujuan untuk menciptakan dasar bagi pencapaian netralitas iklim Uni Eropa, diajukan.

Akhirnya, kerangka kerja Kesepakatan Hijau Eropa merupakan hasil dari pendekatan berlapis yang berfokus pada pemikiran ulang kebijakan untuk penyediaan energi bersih di berbagai bidang. Kesepakatan ini mencakup ekonomi, industri, produksi dan konsumsi, infrastruktur skala besar, transportasi, pangan dan pertanian, konstruksi, perpajakan, dan manfaat sosial. Di semua bidang ini, teknologi digital adalah 'yang krusial untuk mencapai tujuan keberlanjutan dari Kesepakatan Hijau'. Di sini, AI tercantum sebagai salah satu teknologi digital (bersama dengan 5G, komputasi awan dan edge, IoT, dll.), yang akan dieksplorasi karena potensinya untuk mempercepat dan memaksimalkan dampak kebijakan untuk memerangi perubahan iklim dan melindungi lingkungan. Selain teknologi, penekanan diberikan pada data yang dapat diakses dan interoperabel, yang, dikombinasikan dengan infrastruktur digital berbasis superkomputer, komputasi awan, dan jaringan serta solusi kecerdasan buatan, menyediakan kerangka kerja untuk keputusan berbasis bukti dan memperluas kesempatan untuk mengenali, menganalisis, memahami, dan menantang masalah lingkungan.

Namun demikian, seperti halnya dalam kasus emisi, energi terbarukan, dan efisiensi energi, yang telah ditangani oleh UE jauh sebelum abad ke-21, masalah AI terkait energi telah menjadi topik kebijakan Eropa selama beberapa dekade. Aplikasi AI di bidang energi sudah hadir di Eropa. Pada pergantian tahun 1970-an dan 1980-an, lembaga ilmu pengetahuan dan teknologi Komisi – Pusat Penelitian Gabungan (JRC) – memulai studi tentang sistem pakar dan komunikasi manusia-mesin yang dijalankan sebagai bagian dari program penelitian keselamatan nuklir yang digunakan untuk memantau manajemen bahan fisil. Seperti yang diantisipasi, hal ini dapat mengarah pada peningkatan lebih lanjut dalam robotika, yang pada

awal 1980-an Eropa 'cukup unggul dalam hal aspek mikro-mekanika, tetapi agak tertinggal dalam hal robot "cerdas"'. Studi telah menunjukkan bahwa AI dapat diterapkan dalam manajemen industri nuklir, yang mencakup pengambilan keputusan, sistem diagnostik, dan pemodelan perilaku operator. Dua elemen terakhir dibahas lebih lanjut dalam program keselamatan reaktor air ringan (LWR) yang menyediakan penelitian tentang faktor manusia dan interaksi manusia-mesin.

Pada tahun 1990-an, prioritas penelitian dalam kerangka kerja Eropa mengikuti kemungkinan peningkatan keberlanjutan dengan menerapkan teknologi inovatif dan bersih pada sistem produksi. Hal ini menyangkut 'pengembangan 'produksi bersih' dengan, misalnya, menggunakan AI untuk meningkatkan produktivitas, meningkatkan efisiensi energi, atau mengurangi limbah. Contoh tindakan yang diterapkan menunjukkan bahwa hasil ini telah tercapai. Pada tahun 2000-an, selain penelitian, berbagai lembaga Eropa secara bertahap mulai menyadari manfaat penerapan AI di bidang terkait energi. Hal ini berlaku, *antara lain*, untuk situasi di mana robot, otomatisasi canggih, dan AI dapat meningkatkan keselamatan dan efektivitas biaya di tempat kerja.

Pada tahun 2010-an, langkah ini dipercepat, dan berbagai insentif muncul. Mereka termasuk, misalnya rencana UE untuk membuat platform industri digital, yaitu gerbang pasar multi-sisi dari interaksi antara beberapa kelompok pelaku ekonomi. Mereka memiliki beberapa tujuan, termasuk mengatasi tantangan teknologi digital, IoT, data besar, dan sistem cloud otonom, AI, pencetakan 3D, serta investasi dalam inisiatif percontohan dan mercusuar, seperti kota pintar dan lingkungan hidup pintar. Kata sifat - 'pintar' - selama bertahun-tahun, telah menjadi elemen penting dari sistem energi modern. Selain kota pintar dan rumah pintar, orang dapat membuat daftar gagasan seperti manufaktur pintar, mobilitas pintar, pertanian pintar, atau energi pintar, dengan meter pintar dan jaringan pintar. Yang terakhir telah menjadi paradigma pengembangan sektor energi saat ini, di mana data jaringan pintar waktu nyata membantu pengoperasian sistem yang efisien dan memfasilitasi layanan listrik. Namun demikian, dalam semua 'bidang pintar' ini, 'energi' memainkan peran penting. Hal ini mengacu pada optimalisasi penggunaan energi dengan menggunakan berbagai metode dan teknologi untuk menurunkannya, mengadaptasi konsumsi energi berdasarkan tarif dinamis, meningkatkan pengelolaan konsumsi sendiri, penyimpanan dan penyaluran ke jaringan, atau mengubah perilaku dan pola konsumsi.

Meskipun pergeseran teknologi ini, yang melibatkan AI, dapat membantu transisi energi dan membuat langkah (atau bahkan lompatan) menuju ekonomi tanpa emisi, terdapat beberapa kekhawatiran mengenai aspek ekologis dari pergeseran ini. Meningkatnya penggunaan AI mengakibatkan peningkatan penggunaan sumber daya alam, serta meningkatnya permintaan energi dan masalah pembuangan limbah. Meskipun AI, komputasi awan, atau IoT dapat mempercepat dan mengoptimalkan dampak kebijakan terhadap perubahan iklim dan perlindungan lingkungan, pengembangan dan pengoperasiannya harus netral terhadap iklim dan ramah lingkungan. Oleh karena itu, 'pada saat yang sama, Eropa membutuhkan sektor digital yang mengutamakan keberlanjutan'.

10.3 AI DAN KEBIJAKAN IKLIM-ENERGI JEPANG

Sejak tahun 1974, beberapa inisiatif di bidang energi dan lingkungan telah ditetapkan. Bahasa Indonesia: Pada tahun 1974, 'Proyek Sinar Matahari' dimulai, menawarkan kegiatan penelitian dan pengembangan di lima bidang utama: energi surya, energi panas bumi, energi batubara, energi hidrogen, dan penelitian komprehensif. Pada saat krisis energi tahun 1970-an, banyak perhatian diberikan pada energi surya sebagai alternatif baru untuk bahan bakar fosil konvensional. Selain mendorong pertumbuhan produksi sel surya, Proyek Sinar Matahari membentuk dasar untuk kegiatan publik dan swasta lebih lanjut untuk mengembangkan teknologi energi baru. Inisiatif kebijakan pro-lingkungan lainnya dimulai pada tahun 1978. Itu adalah 'Proyek Cahaya Bulan' yang didirikan di Jepang untuk mengarahkan pengembangan teknologi konservasi energi. Program ini melakukan penelitian dan pengembangan bersama antara pemerintah dan sektor industri di bidang teknologi hemat energi, yang terlalu mahal dan berisiko bagi sektor swasta (misalnya turbin gas canggih, teknologi pemanfaatan panas buang, dan pembangkit listrik sel bahan bakar).

Lebih lanjut, sebagai tanggapan terhadap tren stagnasi tahun 1990-an sehubungan dengan penelitian dan pengembangan energi dan untuk memfasilitasi potensi energi terbarukan di Jepang, Proyek Sunshine berkembang menjadi Program Sunshine Baru tahun 1993. Hal ini dimungkinkan oleh reorganisasi, yang mencakup penggabungan 'Proyek Sunshine' dengan 'Proyek Cahaya Bulan' dan 'Program Teknologi Lingkungan Global' pada tahun 1993. Dengan tujuan mengembangkan teknologi inovatif untuk kebutuhan pertumbuhan berkelanjutan, program baru tersebut mencakup pengembangan teknologi fotovoltaik (PV) berbiaya rendah. Bersamaan dengan program insentif seperti 'Program Penyebaran Sistem PV Perumahan', serta pendahulunya 'Program Pemantauan Sistem PV Perumahan', Jepang mampu membangun pasar PV yang mandiri.

Sebagai hasil dari insentif ini, selama tahun 1980-an dan 1990-an, banyak proyek demonstrasi dan penelitian dasar serta pekerjaan pengembangan didukung di Jepang, menciptakan permintaan yang diperlukan untuk sel surya dan peningkatan berkelanjutan dari efisiensi konversi dan ekonomi PV. Berkat proyek-proyek ini, Jepang telah lama menjadi pemimpin dunia dalam energi surya; namun, dengan pengurangan dan kemudian penghentian subsidi surya pada tahun 2005, pasar PV di Jepang telah mandek. Untuk mengubah ini, pada tahun 2009, tujuan nasional untuk meningkatkan kapasitas PV dengan faktor 20 pada tahun 2020 (dengan mempertimbangkan tingkat tahun 2005), dan dengan faktor 40 pada tahun 2040 diumumkan. Pada awal tahun 2010-an, hal ini ditingkatkan dengan tarif feed-in baru untuk produksi listrik dalam instalasi surya. Diskusi lebih lanjut tentang kebijakan energi di Jepang didorong oleh reaktor nuklir Fukushima 2011 kecelakaan.

Pada bulan Juni 2019, Jepang menyetujui 'Strategi Jangka Panjang berdasarkan Perjanjian Paris' (Strategi 2019). Strategi ini membawa visi 'masyarakat terdekarbonisasi' yang berpotensi dicapai oleh Jepang pada paruh kedua abad ke-21 (atau lebih awal) dengan mengurangi emisi gas rumah kaca sebesar 80%. Dengan tujuh bidang utama, banyak operasi terkait iklim akan dilakukan oleh sektor bisnis karena memiliki sumber daya keuangan dan teknologi yang direncanakan Jepang untuk dimanfaatkan dalam aksi iklim negara tersebut.

Dalam hal ini, negara tersebut akan menciptakan lingkungan yang inovatif, mempromosikan dan memanfaatkan 'inovasi untuk dekarbonisasi' yang diterapkan di berbagai tingkatan, dengan partisipasi berbagai pelaku – perusahaan, investor, lembaga keuangan, konsumen, dan pemerintah daerah. Selain elemen bisnis dari inovasi untuk dekarbonisasi, Jepang juga berencana untuk mengatasi masalah sosial yang terkait dengan dekarbonisasi yang diusulkan. Selain tindakan umum seperti memanfaatkan Tujuan Pembangunan Berkelanjutan dan agenda yang ada seperti 'Masyarakat 5.0' ('masyarakat super pintar'), Jepang mengusulkan penciptaan 'Ekonomi Bersirkulasi dan Ekologis'. Hal ini didasarkan pada kerja sama komunitas regional yang memanfaatkan sumber daya mereka secara berkelanjutan untuk menjadi mandiri (semaksimal mungkin), terhubung dalam jaringan komunitas yang bertujuan mencapai dekarbonisasi dan pembangunan berkelanjutan. Dalam kerangka kerja ini, 'komunitas netral karbon' akan memanfaatkan pembangkitan terbarukan dalam jaringan pintar, dan untuk mendorong pemasangan energi terbarukan yang lancar, akan menerapkan, *antara lain*, teknologi blockchain.

Selain itu, model Jepang dari jenis komunitas energi ini juga mempertimbangkan elemen pencegahan bencana, yaitu kemandirian komunitas dengan bantuan jaringan pintar, penyimpanan energi (baterai), sel bahan bakar, serta kogenerasi. Elemen lain dari komunitas ini adalah: menerapkan respons permintaan dan pembangkit listrik virtual yang memperkuat posisi konsumen energi, bersama dengan gagasan 'Sistem Manajemen Energi Desa' (VEMS). Konsep terakhir ini bertujuan untuk mengoptimalkan penggunaan sumber daya energi lokal serta 'pertanian-fotovoltaik', yaitu memasang panel PV di atas lahan pertanian, termasuk yang terbengkalai, dengan cara yang memungkinkan penanaman tanaman. Perlu juga dicatat gagasan untuk menggunakan kombinasi drone, teknologi penginderaan, dan AI untuk mengurangi emisi nitrogen oksida dengan mengendalikan jumlah pupuk yang digunakan serta menggunakan AI untuk memantau emisi gas rumah kaca. Selain itu, Jepang akan mempromosikan penggunaan produk hemat energi baru yang lebih luas dengan menggunakan AI, IoT, dan data besar, yang pasarnya akan tercipta sebelum tahun 2040.

Namun, Strategi 2019 bukanlah referensi pertama yang menyebut AI sebagai solusi potensial bagi Jepang dan sektor energinya. Pada awal 1980-an, Jepang meluncurkan proyek Sistem Komputer Generasi Kelima, sebuah rencana berani untuk merevolusi perangkat keras dan perangkat lunak komputasi. Proyek ini, selain tujuan teknis utamanya, memperkenalkan berbagai tujuan sosial, termasuk peningkatan efisiensi energi. Selain itu, pada tahun 1980-an, aplikasi AI dalam operasi dan pemeliharaan pembangkit listrik tenaga nuklir dipromosikan di bawah program nasional jangka panjang untuk pengembangan dan pemanfaatan energi nuklir di Jepang. Selain itu, pada awal tahun 1980-an, sektor kelistrikan Jepang memulai upaya besar untuk menggunakan alat AI untuk meningkatkan operasi sistem, perencanaan, diagnosis, dan kontrol. Kebutuhan akan solusi ini didorong oleh berbagai alasan, termasuk: kekurangan staf yang berkualifikasi karena pensiunnya pekerja era pasca-Perang Dunia II, keinginan untuk menjalankan sistem dengan margin yang lebih rendah, peningkatan efisiensi, untuk mempercepat pembersihan kesalahan dan pemulihan layanan, meningkatkan hubungan pelanggan, dan menyediakan insinyur utilitas dengan desain dan alat perencanaan yang lebih

efisien. Selain itu, pada saat itu, berbagai platform penelitian digunakan untuk mempromosikan aplikasi AI di Jepang. Masyarakat Riset Koperasi Listrik, Kelompok Litbang untuk AI dalam Utilitas Tenaga Listrik, dan Institut Insinyur Listrik di Jepang (IEEJ) dapat ditemukan di antara mereka.

Namun, pada akhir 1980-an, Jepang mulai kehilangan keunggulannya. Meskipun teknologi yang muncul menjadi lebih dinamis dan lintas disiplin, pemerintah Jepang belum siap untuk memainkan peran konstruktif. Proyek Sistem Komputer Generasi Kelima berakhir pada tahun 1992 tanpa tanda keberhasilan yang jelas. Persaingan birokrasi, koordinasi yang buruk, dan intervensionisme yang kuat menggerogoti posisi menguntungkan perusahaan-perusahaan Jepang dalam perlombaan untuk memimpin bisnis yang digerakkan oleh teknologi informasi. Namun, berbagai jenis penelitian dan pekerjaan di bidang AI di sektor energi masih dilakukan. Contohnya termasuk sistem berbasis pengetahuan eksperimental, Pakar Keamanan Pasokan Energi Jepang (JESSE), yang merupakan model pengambilan keputusan dalam kebijakan energi Jepang. Di luar sistem ini, kegiatan tertentu dilakukan dalam mengembangkan logika fuzzy, jaringan saraf, dan algoritma penalaran mesin untuk aplikasi komersial, juga sebagai bagian dari diskusi tentang prospek kinerja ekonomi Jepang di era pasca-gelembung.

Namun demikian, meskipun memiliki pengalaman luas dalam teknologi, elektronik, dan robotika, beberapa masalah ditemui dalam pertumbuhan industri AI di Jepang. Sebagaimana disorot dalam 'Strategi Teknologi Kecerdasan Buatan' tahun 2017, [k]etika melihat ... makalah yang terkait dengan AI ... jumlah makalah Jepang turun di bawah jumlah makalah di AS dan Tiongkok, dan jelas bahwa tidak ada investasi yang cukup dalam penelitian dan pengembangan oleh sektor publik dan swasta.

Untuk menyediakan lingkungan bagi perkembangan ini, beberapa perbaikan telah dilakukan dalam struktur pemerintahan Jepang. Salah satunya adalah pembentukan Dewan Strategis untuk Teknologi AI. Dewan ini dibentuk pada tahun 2016 dengan tanggung jawab utama untuk mengelola dan mengoordinasikan kegiatan yang dilakukan oleh lima Badan Penelitian dan Pengembangan Nasional, termasuk Organisasi Pengembangan Energi dan Teknologi Industri Baru (NEDO). NEDO, yang berfokus pada sistem energi, konservasi energi, dan lingkungan, serta teknologi industri dengan kemajuan dalam robotika, AI, dan IoT, mendanai proyek-proyek penghematan energi dan permintaan energi yang menerapkan atau menyelidiki berbagai solusi AI.

Akhirnya, isu-isu manajemen energi, pasokan energi, konsumsi energi, dan penghematan energi telah ditangani, antara lain, dalam 'Strategi AI 2019' Jepang baru-baru ini. Strategi tersebut mencakup insentif seperti pengembangan sistem energi yang mandiri dan terdistribusi, tahan bencana, pengumpulan dan peninjauan data penggunaan energi, penyampaian pesan yang disesuaikan untuk setiap pengguna, promosi tindakan penghematan energi melalui kombinasi wawasan perilaku seperti nudge dan boost (dengan teknologi canggih seperti AI dan IoT), atau penyediaan infrastruktur dan platform untuk data besar energi. Selain itu, jenis solusi AI ini dapat membantu mencapai tujuan mengatasi kendala energi lingkungan dengan memanfaatkan teknologi mutakhir Jepang.

10.4 MODEL REGULASI CERDAS

Selain implikasi kebijakan AI terhadap isu iklim dan energi, AI juga berdampak pada legislasi. Pada tahun 2016, Panel Penilaian Opsi Sains dan Teknologi (STOA) Parlemen Eropa mempresentasikan laporan tentang undang-undang Uni Eropa yang akan terpengaruh oleh perkembangan di bidang robotika dan AI, salah satunya adalah energi. Area yang menjadi perhatian terkait energi antara lain mencakup potensi penyalahgunaan atau penangkapan infrastruktur robotika yang diciptakan untuk pasokan listrik, atau peninjauan pelabelan (efisiensi energi, desain ramah lingkungan, dan informasi produk standar). Selain itu, beberapa instrumen hukum yang mungkin perlu ditinjau atau diperbarui juga dicantumkan, misalnya, Arahan Efisiensi Energi atau Arahan Pelabelan Energi. Bidang lain yang dapat memperoleh manfaat dari penerapan AI adalah regulasi di sektor kelistrikan. Sementara literatur menyajikan banyak definisi dan cara berbeda untuk menafsirkan regulasi, dan cara regulasi ditafsirkan bervariasi di antara bahasa dan budaya hukum, ketika membahas regulasi, orang mungkin menemukan bahwa regulasi memiliki aplikasi yang sempit dan luas. Aplikasi yang sempit mengacu pada seperangkat aturan yang diadopsi berdasarkan undang-undang yang mengikat, sementara aplikasi yang luas dapat didefinisikan oleh mekanisme kontrol dan pengaruh sosial apa pun. Selain itu, regulasi dalam aplikasinya yang sempit juga dapat didefinisikan dengan merujuk pada aktivitas otoritas regulasi, di mana regulasi adalah regulator, yang ditunjuk untuk mengatasi kekurangan pasar dalam memenuhi tujuan kepentingan publik (atau kolektif). Selain lingkungan, iklim, dan keberlanjutan, yang disajikan dalam bab ini, tujuan-tujuan ini mungkin (atau secara tradisional) berorientasi pada persaingan dan keterjangkauan (pasar energi dan konsumen energi), serta keamanan energi. Seperti yang dibahas dalam bab ini, AI dapat – dan memang – membantu mencapai tujuan-tujuan ini.

Namun demikian, yang tersisa untuk analisis lebih lanjut adalah aktivitas regulator energi yang menggunakan AI. Pada kenyataannya, terlepas dari banyak contoh kemajuan teknologi dalam AI, administrasi publik (termasuk otoritas regulator) masih menyediakan layanan dengan cara lama, meskipun penerapan teknologi dan solusi baru dalam administrasi publik dapat sangat bermanfaat. Hal ini dapat meningkatkan efektivitas dan efisiensi pemerintah, meningkatkan kepuasan warga, memfasilitasi konektivitas dan interaktivitas, serta mengurangi beban administratif. Perbaikan semacam itu di sektor kelistrikan dapat membantu memfasilitasi perubahan yang lebih besar jauh lebih cepat daripada yang diperkenalkan oleh administrasi yang bertanggung jawab atas energi saja.

Sejarah menunjukkan bahwa peraturan tertentu (atau ketiadaan peraturan) dapat menyebabkan masalah. Contoh yang baik adalah masalah deregulasi yang terjadi di California pada tahun 2000-an, yang memengaruhi kualitas layanan, serta harga. Hal ini juga dapat terjadi di belahan dunia lain, termasuk Jepang, jika reformasi pasar energi gagal. Untuk menghindari hal ini, solusinya bisa jadi adalah apa yang disebut model regulator 'penjaga harian'. Penjaga harian menjembatani kesenjangan antara subordinasi total dan pelepasan penuh. Di bawah model ini, negara tidak berpartisipasi aktif di pasar, juga tidak sepenuhnya absen (keseimbangan antara pasar dan negara, dan antara kepentingan publik dan swasta).

Untuk memberikan keseimbangan ini, regulator adalah pengamat yang waspada, ingin tahu, dan aktif (tetapi bukan peserta aktif). Regulator adalah penjaga harian yang, bertindak di bawah kekuasaan negara, dapat mengklarifikasi aturan pasar, memberikan informasi kepada pelaku pasar, dan menegakkan aturan melalui penggunaan sanksi.

Namun, dalam kenyataannya model ini dapat terbukti tidak memadai karena regulator hanyalah manusia, dan orang yang salah di posisi yang salah dapat menyebabkan regulasi yang tidak memadai dan tidak efektif. Setiap solusi yang baik membutuhkan pintu air, misalnya badan-badan kolegial dalam sistem regulasi (komisi, dewan, dll.), yang dapat mengoreksi kesalahan regulator. Namun demikian, tidak seorang pun dapat sepenuhnya menghindari penunjukan tiga, lima, atau tujuh (atau jumlah lainnya) kandidat yang tidak sesuai sebagai komisarisi komisi regulasi tersebut, sebagaimana ditunjukkan oleh pengalaman sebelumnya.

Jika menyangkut alternatif, salah satu solusinya adalah kemungkinan penerapan AI untuk mendukung model day-watchman. Teknologi dan solusi baru dapat membuat model ini berfungsi dengan membuatnya cukup cerdas untuk mengatur sektor kelistrikan yang cerdas. Namun, hal ini menimbulkan masalah lain. Jika alternatif ini layak, bagaimana model ini dapat dikembangkan dalam sistem hukum? Bagaimana AI dapat meningkatkan perangkat regulasi tradisional di sektor kelistrikan? Dapatkah AI diperlakukan sebagai perangkat terpisah atau titik validasi regulasi? Jenis risiko apa yang mungkin muncul? Dapatkah risiko ini dimitigasi, dan jika demikian, apa yang diperlukan untuk memitigasinya? Pertanyaan-pertanyaan ini memerlukan penelitian lebih lanjut tentang regulator energi dan penerapan AI, baik di sektor kelistrikan maupun dalam tugas regulasi mereka.

10.5 KESIMPULAN

Pada bulan Oktober 2020, Dewan Eropa memutuskan untuk menyediakan setidaknya 20% dana di bawah Fasilitas Pemulihan dan Ketahanan untuk transisi digital. Dana ini, bersama dengan dana di bawah MFF, diharapkan dapat membantu memajukan tujuan lingkungan dan iklim Uni Eropa dengan 'memanfaatkan potensi penuh teknologi digital'. Selain itu, Area Penelitian Eropa (ERA) yang baru akan meningkatkan pemulihan Eropa dan membantu transformasi hijau dan digitalnya dengan mendorong daya saing berbasis inovasi dan mempromosikan kedaulatan teknologi di bidang-bidang strategis utama seperti AI, data, mikroelektronika, komputasi kuantum, 5G, penyimpanan energi, energi terbarukan, hidrogen, nol emisi, dan mobilitas cerdas.

Selain menjadi elemen pendekatan strategis di bidang energi dan lingkungan, yang dapat dilihat dalam dokumen-dokumen yang diadopsi, seperti strategi, kebijakan, atau program, AI menjadi fokus tindakan legislatif. Meskipun Arahan Efisiensi Energi dan Arahan Pelabelan Energi telah di amandemen atau dicabut, amandemen yang diadopsi tidak menyebutkan AI. Namun, AI disebutkan di bawah rezim Uni Eropa tentang desain ramah lingkungan. Peraturan Komisi (UE) 2019/424,, yang menyediakan spesifikasi desain ramah lingkungan untuk server dan produk penyimpanan data, memperkenalkan definisi 'server Komputasi Kinerja Tinggi (HPC)', yang membahas AI dan menetapkan aturan tentang efisiensi jenis server ini. Sebagaimana diungkapkan dalam pembukaan Peraturan ini,

persyaratan e-codesign harus menyelaraskan persyaratan konsumsi energi dan efisiensi sumber daya untuk server dan produk penyimpanan data di seluruh Uni, agar pasar internal dapat beroperasi lebih baik dan untuk meningkatkan kinerja lingkungan dari produk-produk tersebut.

Dalam konteks ini, tidak boleh diabaikan bahwa isu-isu terkait konsumsi energi merupakan bagian permanen dari diskusi tentang penggunaan AI di Uni Eropa. Didorong oleh aksi iklim, diarahkan oleh pertumbuhan sumber energi terbarukan, pengurangan emisi, dan peningkatan efisiensi energi, kebijakan energi Eropa bergerak menuju netralitas iklim, yang ditetapkan akan tercapai pada tahun 2050. Di Jepang, pemerintahan Perdana Menteri Suga berjanji untuk mengurangi emisi gas rumah kaca di Jepang menjadi nol bersih pada tahun yang sama.

Namun, transisi ini tidak akan terjadi dengan sendirinya, dan masih banyak upaya yang diperlukan. Misalnya, dengan mendirikan badan baru pada tahun 2021, Jepang berencana untuk meningkatkan digitalisasi fungsi pemerintahan. AI, robotika, IoT, komputasi awan, data besar, dan banyak inovasi teknologi lainnya memiliki potensi besar untuk meningkatkan proses ini, menjadikan transisi cerdas, berkelanjutan, hijau, dan efisien. Insentif yang serupa dengan Green Deal atau aksi pascapandemi COVID-19 merupakan kerangka kerja yang tepat tidak hanya untuk membuka kapasitas AI dalam transformasi ini, tetapi juga untuk mengelola dan mengaturnya secara cerdas. Hal ini juga berlaku untuk pandemi COVID-19 yang sedang berlangsung, di mana AI menemukan beragam aplikasi untuk mengelola perubahan permintaan listrik (peningkatan perumahan – penurunan industri/komersial) akibat pergeseran pola kerja dan gaya hidup. Insentif tersebut mencakup, *antara lain*, pemantauan produksi dan konsumsi energi secara real-time dari berbagai sumber, termasuk energi terbarukan, analitik data, dan pengembangan metode prediktif untuk menganalisis tren penggunaan energi yang fluktuatif.

Dalam hal ini, AI juga harus diakui secara jelas sebagai alat untuk mencapai netralitas iklim, sehingga pengembangannya harus rendah emisi, didorong oleh energi terbarukan, dan hemat energi. Hal ini juga menyangkut usulan regulasi terkini yang menetapkan aturan harmonisasi tentang AI (Undang-Undang Kecerdasan Buatan) di UE. Meskipun membahas AI sebagai keunggulan kompetitif utama untuk mendukung hasil yang bermanfaat secara sosial dan lingkungan, di bidang-bidang seperti efisiensi energi dan mitigasi serta adaptasi perubahan iklim, dan secara singkat membahas keberlanjutan lingkungan, Undang-Undang Kecerdasan Buatan yang diusulkan tidak mencakup ketentuan khusus tentang peningkatan isu iklim dan energi melalui penggunaan AI. Jika kita benar-benar ingin mempersiapkan diri untuk realitas masa depan, kita harus bertindak sekarang dan cepat, karena kemarin telah menjadi hari ini di banyak bidang yang terkait dengan AI.

BAB 11

REGULASI AI MILITER UNTUK PERLINDUNGAN SIPIL

11.1 PENDAHULUAN

Menurut Mahkamah Internasional, prinsip-prinsip hukum konflik bersenjata (LOAC) berlaku 'untuk semua bentuk peperangan dan semua jenis senjata, baik di masa lalu, masa kini, maupun masa depan.' Salah satu diskusi yang paling sulit dan kompleks tentang penerapan hukum internasional pada teknologi baru adalah pada sistem senjata otonom (AWS). Selama bertahun-tahun, diskusi ini telah bergulir dari halaman-halaman jurnal akademis dan aula-aula Palais des Nations di Jenewa, tempat Kelompok Pakar Pemerintah tentang Sistem Senjata Otonom Mematikan (GGE on LAWS) rutin bersidang. Kelompok ini tidak hanya berkontribusi besar pada kecanggihan perdebatan tentang senjata-senjata ini, tetapi juga memfasilitasi pertukaran posisi Negara yang dinamis mengenai ruang lingkup aturan hukum internasional dan penerapannya pada kemampuan militer yang muncul dengan berbagai tingkat otonomi. Kelompok ini telah mengadopsi, secara konsensus, sebelas Prinsip Panduan, yang pertama menegaskan bahwa '[i]hukum humaniter internasional terus berlaku sepenuhnya untuk semua sistem persenjataan, termasuk potensi pengembangan dan penggunaan AWS yang mematikan.' Meskipun terdapat kesepakatan bahwa hukum humaniter internasional berlaku untuk semua sistem persenjataan, pertanyaan sulit yang masih tersisa adalah bagaimana tepatnya hukum tersebut berlaku untuk sistem-sistem tersebut. Pertanyaan '*bagaimana*' ini terkait dengan ruang lingkup aturan hukum humaniter internasional, termasuk kontur tepat tugas-tugas yang dibebankan kepada pihak-pihak dalam konflik bersenjata.

Tujuan Bab ini adalah untuk memberikan gambaran umum mengenai posisi Negara dalam pengaturan kemampuan militer otonom yang mengandalkan AI, dan untuk menguraikan tiga kewajiban yang sangat penting dalam memastikan penyebaran senjata yang aman dengan otonomi dalam pengambilan keputusan: kewajiban untuk melakukan tinjauan hukum terhadap senjata baru, kewajiban untuk mengambil tindakan pencegahan dalam serangan, dan kewajiban untuk mengambil tindakan pencegahan terhadap dampak serangan.

11.2 OTONOMI DAN SARANA PERANG

Untuk mencapai tujuan melemahkan potensi militer musuh, pihak yang bertikai menggunakan sarana perang. Kategori ini mencakup 'senjata, sistem persenjataan, atau platform yang digunakan untuk tujuan serangan dalam konflik bersenjata'. Berdasarkan Hukum Humaniter Internasional, 'hak Para Pihak yang berkonflik untuk memilih [...] sarana perang tidaklah tak terbatas', dan pembatasan atas hak ini ada dalam bentuk (1) larangan umum terkait dengan kelas-kelas bahaya – sarana untuk menyebabkan cedera yang berlebihan atau penderitaan yang tidak perlu; sarana yang dimaksudkan, atau dapat diperkirakan, untuk menyebabkan kerusakan yang meluas, jangka panjang, dan parah terhadap lingkungan alam; sarana yang tidak dapat diarahkan pada tujuan militer tertentu; sarana yang dampaknya tidak

dapat dibatasi, dan (2) larangan khusus terhadap jenis senjata tertentu, seperti larangan ranjau anti-personel dan senjata laser yang dirancang khusus untuk menyebabkan kebutaan permanen, antara lain.

Wajar saja, negara-negara memiliki kepentingan dalam mengembangkan alat perang yang andal dan efektif. Di antara banyak pertimbangan yang mendorong perkembangan teknologi di bidang persenjataan adalah keinginan untuk mengatasi kelemahan manusia, melindungi pasukan sendiri, dan meminimalkan biaya. Otonomi, dalam beberapa tahun terakhir, telah dipandang sebagai fitur senjata yang menarik, terutama karena menjanjikan tingkat presisi tinggi yang dicapai dari jarak aman. Oleh karena itu, tidak mengherankan bahwa banyak negara telah membuat kemajuan signifikan dalam penelitian dan pengembangan senjata dan sistem persenjataan tersebut, dan bahwa kemajuan ini hanya menandai awal dari otomatisasi medan perang. Dalam laporan yang baru-baru ini diperbarui yang disiapkan untuk Anggota dan Komite Kongres Amerika Serikat yang berjudul 'Kecerdasan Buatan dan Keamanan Nasional', Kongres AS mengonfirmasi bahwa AS sudah mengintegrasikan AI ke dalam pertempuran melalui Proyek Maven, dan tidak menganggap bahwa LOAC memberlakukan larangan pengembangan AWS yang mematikan. Rusia juga telah menunjukkan komitmennya untuk memanfaatkan AI di seluruh sektor melalui Strategi AI 2019. Di bidang militer, perusahaan Kalashnikov, anak perusahaan dari perusahaan milik negara Rusia Rostec, mengakui penelitian dan pengembangan senjata berbasis AI saat ini, dan Presiden Vladimir Putin, dalam pidatonya tahun 2018 di Majelis Federal, mengumumkan bahwa Rusia telah mengembangkan kendaraan selam tak berawak yang dapat membawa hulu ledak konvensional atau nuklir. Tiongkok juga berupaya memodernisasi militernya dengan mengoperasionalkan AI.

Namun, apa sebenarnya otonomi yang ingin dicapai oleh negara-negara ini dan negara-negara lain? Otonomi dalam persenjataan dapat terwujud dalam berbagai cara: otonomi dalam pengumpulan informasi, otonomi dalam mengidentifikasi target, otonomi dalam mengerahkan kekuatan terhadap seseorang atau objek. Dari jenis-jenis otonomi ini, yang dianggap paling bermasalah adalah otonomi yang bergantung pada rangkaian sensor dan algoritma komputer untuk secara independen mengidentifikasi target dan mengerahkan kekuatan terhadap target tersebut tanpa kendali manusia manual dan pengambilan keputusan konkret manusia untuk menyerang target tertentu.

AI kemungkinan akan memainkan peran kunci dalam pengembangan otonomi semacam itu untuk menyerang langsung orang dan objek. Hal ini karena AI dipandang sangat cocok untuk memenuhi kebutuhan medan perang modern: sebuah algoritma yang dapat belajar dari lingkungannya dan beradaptasi dengan perubahan keadaan bahkan di lingkungan dengan komunikasi yang buruk atau terputus, dan yang selanjutnya dapat memungkinkan serangan presisi di area yang mungkin tidak dapat diakses oleh pasukan darat. Khususnya, kemajuan dalam pembelajaran mesin telah memungkinkan untuk membayangkan suatu entitas yang mampu belajar, menjadi lebih baik melalui pengalaman, melalui perbandingan statistik dalam data masa lalu. Aplikasi militer pembelajaran semacam itu akan memungkinkan negara-negara untuk menghindari pemrograman manual yang terperinci, yang dapat

memakan waktu dan membutuhkan banyak penelitian. Boothby memberikan contoh sebuah mesin yang 'dapat mengamati pola kehidupan di area yang diminati dan kemudian dapat menyesuaikan daftar target yang dimasukkan ke dalamnya sebelum misi untuk memperhitungkan pengamatan yang telah dilakukannya'.

Bahwa pengenalan AI dalam pengambilan keputusan tentang keterlibatan langsung target akan memerlukan perubahan dalam prosedur pelatihan, tinjauan hukum, dan pengerahan tidak dapat disangkal. Namun, yang diperdebatkan *adalah* apakah pengenalan kemampuan AI tersebut akan mengarah pada penilaian ulang yang mendasar terhadap cara senjata dikonseptualisasikan dan diatur. Untuk mempertimbangkan pertanyaan ini, ada baiknya membandingkan aplikasi militer baru ini dengan senjata sebelumnya. Studi ini berfokus pada ranjau darat dan peluru kendali sebagai contoh.

Ranjau darat menghasilkan ledakan setelah mengarahkan gaya kinetik tertentu. Setelah ditempatkan di lokasi tertentu, tidak perlu ada keputusan mengenai dampak spesifiknya. Dalam hal ini, ranjau darat memiliki tingkat otonomi tertentu dibandingkan mereka yang menggunakannya untuk membunuh, melukai, merusak, atau menghancurkan pasukan atau material musuh. Justru karena alasan itulah, ranjau darat anti-personel telah dilarang oleh Konvensi Larangan Penggunaan, Penimbunan, Produksi, dan Pemindahan Ranjau Anti-Personel serta Pemusnahannya. Namun, kurangnya pengambilan keputusan konkret mengenai dampak spesifik oleh angkatan bersenjata tidak berarti bahwa ranjau darat membuat 'keputusan' sendiri. Kepastian, kondisi untuk terlibat adalah karakteristik utama senjata ini; kondisi yang tepat untuk memicu ledakan telah diketahui secara pasti sebelumnya.

Rudal berpemandu dapat dikerahkan melintasi jarak geografis yang luas, bahkan tanpa kedekatan fisik. Yang menjadi ciri senjata-senjata ini adalah, pertama, fakta bahwa mereka *dipandu* melalui jalurnya, dan kedua fakta bahwa tidak ada independensi dalam pemilihan target. Namun, jarak geografis dan selang waktu antara peluncuran dan dampak dapat memengaruhi beberapa faktor intervensi, termasuk perubahan mendadak dalam kondisi cuaca. Seperti ranjau darat, ada pemahaman tertentu tentang faktor-faktor yang dapat memengaruhi kinerja senjata-senjata ini, termasuk hubungannya dengan lingkungan.

Lalu, apa yang membuat kita tidak nyaman dengan pengerahan senjata yang dapat terlibat dalam penargetan langsung orang dan benda tanpa keputusan penargetan khusus yang dibuat oleh manusia? Ada dua faktor yang dapat menjelaskan kegelisahan ini: yang pertama adalah penerimaan bahwa senjata-senjata ini dapat belajar, sehingga bertindak dengan cara yang tidak diprediksi secara spesifik sebelum keterlibatan target, dan yang kedua adalah potensi pengambilan keputusan AI yang tidak dapat dijelaskan. Kedua faktor ini terkait erat.

Terkait poin pertama, semua senjata memiliki tingkat ketidakpastian tertentu dalam pertempuran konkret. Misalnya, senapan dapat meleset, dan sinar laser dapat dibelokkan karena kondisi cuaca. Meskipun demikian, kondisi cuaca buruk dapat memengaruhi sinar laser, meskipun kita tidak dapat sepenuhnya memprediksi perubahan cepat dalam kondisi cuaca untuk serangan tertentu. Dengan senjata AI, mungkin ada titik di mana akan diterima bahwa metode serangan target senjata tidak perlu sepenuhnya dapat diprediksi oleh personel

bersenjata yang mengerahkannya jika tetap dapat diprediksi bahwa metode tersebut akan tetap berada dalam parameter terprogram tertentu yang dapat diterima. Jarak spasial dan temporal tidak, dengan demikian, menentukan peningkatan risiko ketidakpastian. Seseorang dapat mengerahkan rudal jarak jauh dengan prediksi penuh tentang keadaan yang dapat memengaruhi lintasannya. Namun, seperti halnya senjata tradisional, jarak spasial dan temporal dapat meningkatkan risiko faktor-faktor yang mengganggu pengoperasian senjata.

Terkait poin kedua, kemungkinan menerima tingkat ketidakpastian dalam metode operasi membawa risiko tambahan berupa 'ketidakjelasan'. Kurangnya pemahaman menyeluruh tentang cara senjata AI berinteraksi dengan lingkungannya dapat menghambat deteksi tanda-tanda peringatan oleh mereka yang mengoperasikannya, serta upaya selanjutnya untuk memahami alasan di balik keterlibatan target tertentu yang tidak disengaja. Kebutuhan untuk memastikan 'pemahaman menyeluruh tentang kemampuan dan keterbatasan senjata' telah ditegaskan oleh delegasi negara.

Dengan semakin banyaknya tugas dan pengambilan keputusan yang dialihkan ke senjata, perbedaan antara alat perang dan agen manusia menjadi lebih sulit dipertahankan. Hal ini karena senjata secara tradisional dipahami sebagai 'instrumen ekstrakorporeal', 'perangkat, sistem, amunisi, alat, substansi, objek, atau peralatan', 'mekanisme untuk mencapai tujuan'. Intinya, hal ini tampaknya menyiratkan tidak lebih dari sekadar instrumen untuk menyalurkan kehendak manusia. Meskipun AI dipengaruhi oleh kemauan tim pemrogram, penasihat hukum, dan komandan yang menentukan tujuan, tugas, dan parameter senjata, ruang pengambilan keputusan independen yang muncul dapat terlihat. Dalam ruang yang muncul ini, dapat terdapat keputusan tentang pola penerbangan optimal menuju target tertentu, keputusan yang dikirimkan kepada operator manusia mengenai identifikasi target, atau bahkan keputusan independen untuk secara langsung menyerang orang atau objek tertentu. Menurut Liu, 'kerangka kerja otonomi mengedepankan status liminal AWS antara agen dan objek', dan kesamaan berbasis otonomi ini bahkan mendorong beberapa orang untuk menyebut 'robot tempur' dan 'agen perangkat lunak masa perang'.

Apa pun kesamaan faktual yang mungkin ada antara kekuatan pengambilan keputusan yang diberikan kepada senjata berbasis AI dan yang biasanya dimiliki oleh agen manusia, garis normatif di bawah LOAC tetap tidak berubah. Ini karena negara secara konsisten merujuk pada aplikasi militer otonom sebagai senjata, dan menegaskan bahwa LOAC 'memberlakukan kewajiban pada negara, pihak dalam konflik bersenjata dan individu, bukan mesin'. Tentu saja ada alasan yang sangat bagus untuk mempertahankan garis ini, mungkin yang paling penting yang berakar dalam memastikan bahwa tanggung jawab tidak akan dialihkan dari negara dan komandan manusia melalui pengenalan agensi mesin. Namun, penting untuk diingat bahwa pilihan paradigma hukum antara senjata dan agen tidak bebas konsekuensi. Ketika anggota angkatan bersenjata melanggar aturan penargetan dengan cara yang, dari perspektif rantai komando resmi, tidak dapat diprediksi (misalnya, tindakan penargetan yang bertentangan dengan instruksi), tindakan mereka akan tetap menimbulkan tanggung jawab negara. Menurut Pasal. 91 Protokol Tambahan I, pihak yang berkonflik 'bertanggung jawab atas *semua tindakan* yang dilakukan oleh orang-orang yang menjadi

bagian dari angkatan bersenjata'. Namun, gambaran yang muncul dari senjata yang 'bertindak' secara tak terduga jauh lebih kompleks. Ketidakpastian dapat menjadi indikator bahwa senjata tersebut pada dasarnya tidak pandang bulu, dan oleh karena itu dilarang, sebagai *suatu jenis senjata*, berdasarkan LOAC. Lebih lanjut, suatu senjata terkadang dapat 'bertindak' dengan cara yang tidak terduga bagi operatornya dan tanpa secara inheren tidak pandang bulu. Tampaknya telah ditetapkan dengan baik bahwa rudal dapat menyimpang dari jalurnya dan peluru dapat memantul, dan bahwa kecelakaan semacam itu tidak akan tercakup dalam aturan-aturan pelarangan LOAC. AWS, seperti semua senjata, akan rentan terhadap malfungsi. Misalnya, larangan mengarahkan serangan terhadap warga sipil, dapat dibayangkan bahwa seorang komandan dapat bertujuan untuk mengarahkan serangan terhadap sasaran militer yang sah, namun, melalui malfungsi, tindakan kekerasan yang dimediasi melalui senjata tersebut akan mengenai sasaran yang berbeda dari yang diinginkan oleh komandan. Konsekuensi dari serangan tersebut tidak akan, dengan demikian, mengubah kualifikasi tindakan pengarah awal. Dalam kasus seperti itu, untuk menentukan apakah telah terjadi pelanggaran Hukum Humaniter Internasional, perlu untuk memeriksa kewajaran pengarah kemampuan khusus ini, dengan mempertimbangkan semua risiko, dalam keadaan serangan tersebut.

Berdasarkan hal di atas, tingkat ketidakpastian tertentu dapat menunjukkan bahwa suatu senjata dilarang *secara khusus*, karena tidak mampu diarahkan pada sasaran militer tertentu. Pada saat yang sama, semua senjata memiliki tingkat ketidakpastian malfungsi tertentu, dan apa yang tidak dapat diprediksi dalam penempatan tertentu (misalnya, pantulan peluru), pada kenyataannya, merupakan fitur senjata yang dapat diprediksi. Apakah suatu senjata secara inheren dilarang berdasarkan LOAC termasuk dalam pertanyaan yang diajukan oleh 'hukum senjata'. Namun, kesulitan nyata dalam menentukan apakah penggunaan senjata berbasis AI tertentu telah dilakukan dengan melanggar hukum humaniter akan terletak pada ketidakpastian antara senjata yang tidak pandang bulu dan senjata yang beroperasi dalam batas-batas yang wajar dan dapat ditoleransi dari malfungsi yang diharapkan.

Kecuali jika dapat ditunjukkan bahwa ada sesuatu tentang senjata yang mengandalkan AI untuk penargetan yang menempatkannya, sebagai suatu kategori, dalam kotak regulasi tertentu di bawah Hukum Humaniter Internasional, analisisnya harus lebih rinci dan berdasarkan karakteristik model senjata AI tertentu. Sama seperti dalam kategori bom pesawat yang lebih luas, yang pada dasarnya *tidak sembarangan*, kita memiliki bom udara modifikasi yang *tidak* pandang bulu. Oleh karena itu, penting untuk menilai setiap sistem senjata berdasarkan keunggulannya sendiri. Dan bahkan jika AWS tertentu tidak secara inheren sembarangan, mungkin saja senjata semacam itu hanya dapat digunakan di wilayah non-perkotaan. Senjata berbasis AI tidak akan muncul sebagai kelompok monolitik. Di masa mendatang, pertimbangan cermat terhadap berbagai jenis kemampuan otonom dan tujuan penggunaannya akan menjadi faktor krusial untuk memastikan adanya batasan yang berarti dalam penerapan, termasuk dalam klarifikasi lebih lanjut mengenai jaminan prosedural untuk mitigasi risiko.

11.3 SENJATA OTONOM DAN PERDEBATAN GGE TENTANG LAWS

Dalam beberapa tahun terakhir, perdebatan tentang regulasi AWS telah mencapai tingkat kecanggihan yang lebih tinggi. Hal ini sebagian besar disebabkan oleh pertemuan Kelompok Pakar Pemerintah (GGE) tentang LAWS. Kelompok ini menciptakan platform untuk pertukaran pandangan yang dinamis antara negara dan aktor non-negara, termasuk melalui makalah kerja yang diserahkan kepada kelompok tersebut. Hasilnya adalah 11 Prinsip Panduan yang diadopsi melalui konsensus dan menguraikan faktor penentu kesepakatan. Pada saat yang sama, delegasi terus berbeda pendapat dalam hal-hal tertentu, dan perbedaan ini disorot dalam laporan yang diadopsi oleh GGE. Meskipun Prinsip Panduan penting sebagai penegasan titik temu, prinsip-prinsip tersebut tidak menyelesaikan banyak kontroversi tajam dan tidak menghalangi pemahaman yang berbeda tentang isi LOAC terkait AWS. Misalnya, sementara delegasi sepakat tentang peran kunci unsur manusia dalam penggunaan senjata tersebut, tidak ada konsensus tentang kontur unsur tersebut, aplikasi konkretnya, dan struktur 'interaksi manusia-mesin'. Topik yang sangat sulit adalah tentang tingkat keterlibatan manusia yang diperlukan, tempatnya dalam rantai operasional, dan apakah masukan manusia pada tahap pengembangan (pemrograman) akan cukup untuk memungkinkan kepatuhan terhadap Hukum Humaniter Internasional. Selain itu, ada ketidaksepakatan yang signifikan tentang kemungkinan memastikan penargetan yang sesuai dengan Hukum Humaniter Internasional (HHI) tanpa *masukan* manusia langsung pada *target tertentu*. Beberapa delegasi telah mempertimbangkan tindakan pencegahan - seperti pengujian, pelatihan, dan penetapan prosedur - sebagai cara yang memadai untuk memastikan pengoperasian senjata yang konsisten dengan aturan HHI; yang lain telah mengajukan keraguan serius tentang kemungkinan beroperasi di lingkungan operasional yang kompleks tanpa 'penilaian manusia dan penilaian berbasis konteks'.

Pada tingkat umum yang sangat tinggi, dapat dikatakan bahwa semua pihak di GGE tentang LAWS sepakat bahwa AWS memiliki hubungan khusus dengan risiko medan perang. Namun, sementara beberapa delegasi melihat senjata ini sebagai cara untuk menghindari risiko tradisional, yang lain menganggap senjata itu sendiri sebagai risiko yang signifikan dan belum pernah terjadi sebelumnya bagi warga sipil. Misalnya, dalam laporan tahun 2020 'Kesamaan dalam Komentar Nasional tentang Prinsip-Prinsip Panduan', Ketua Kelompok Pakar Pemerintah saat itu, Yang Mulia Bapak Jānis Kārkliņš, mencatat bahwa 'beberapa komentar berpendapat bahwa teknologi yang muncul di bidang AWS yang mematikan dapat mendukung implementasi Hukum Humaniter Internasional karena, *antara lain*, pengurangan kesalahan dan risiko yang berhubungan dengan manusia, peningkatan presisi dan akurasi, dan kemampuan untuk menggabungkan mekanisme penghancuran diri, penonaktifan diri, atau netralisasi diri. Sementara yang lain berpendapat bahwa hasil ini tidak dapat dipastikan dan tidak boleh diasumsikan.

Lebih lanjut, Komite Internasional Palang Merah (ICRC) dan Institut Penelitian Perdamaian Internasional Stockholm (SIPRI) menerbitkan laporan penting tentang konsep kendali manusia yang bermakna, yang kemudian dikutip oleh delegasi negara. Laporan tersebut mengeksplorasi tingkat kendali manusia yang perlu dipastikan dalam rangka

memitigasi risiko yang ditimbulkan oleh AWS. Untuk memitigasi ketidakpastian yang melekat, laporan tersebut mempertimbangkan pentingnya kendali atas parameter penggunaan sistem persenjataan (misalnya, keberadaan mekanisme pengaman), kendali atas lingkungan (misalnya, kendala temporal dan spasial dalam pengerahan), dan kendali melalui interaksi manusia-mesin (misalnya, pengawasan manusia dan kapasitas untuk campur tangan dalam pengoperasian senjata).

Kekhawatiran atas tingkat risiko yang tidak dapat diterima merupakan inti dari diskusi ini. Risiko ini, yang terwujud melalui pengenalan elemen lain dari ketidakpastian dan ketidakpastian medan perang, dapat menimbulkan bahaya baru bagi orang dan objek yang dilindungi. Beberapa delegasi di GGE menganggap bahwa hukum internasional, dan khususnya LOAC, memberikan kendala yang signifikan terhadap pengembangan dan pengerahan AWS. Pandangan ini telah dijabarkan dalam berbagai pernyataan dan makalah kerja yang diajukan oleh peserta debat. Bagian berikut akan berfokus pada tiga kewajiban perlindungan yang timbul berdasarkan LOAC, dan akan mengkaji jaminan yang diberikannya dalam pengembangan dan penerapan AWS.

11.4 KEWAJIBAN PERLINDUNGAN

Berdasarkan Pasal 51(1) Protokol Tambahan I, 'penduduk sipil dan individu sipil harus menikmati perlindungan umum terhadap bahaya yang timbul dari operasi militer' dan berdasarkan Pasal 57(1), 'dalam pelaksanaan operasi militer, perhatian terus-menerus harus diberikan untuk melindungi penduduk sipil, warga sipil, dan objek sipil'. Warga sipil terpapar berbagai ancaman pada saat konflik: diteror secara sengaja atau digunakan sebagai tameng manusia, atau menjadi korban tambahan dalam konteks serangan terhadap sasaran militer yang sah, di antara yang lainnya. Bahaya yang dihadapi warga sipil tidak hanya luas, tetapi juga terus berkembang. Urbanisasi konflik merupakan contoh nyata dari perubahan lanskap bahaya, demikian pula kemajuan teknologi di bidang persenjataan. Dalam konteks ini, kewajiban tertentu yang bersifat prosedural dapat memperoleh signifikansi khusus, karena kewajiban tersebut berupaya menjamin penghapusan atau setidaknya minimalisasi risiko medan perang. Mengakui pentingnya langkah-langkah tersebut, negara-negara yang berpartisipasi dalam GGE tentang LAWS sepakat secara konsensus dengan Prinsip Panduan (g), yang menyatakan bahwa '[p]enaksiran [r]isiko dan langkah-langkah mitigasi harus menjadi bagian dari siklus perancangan, pengembangan, pengujian, dan penyebaran teknologi baru dalam sistem persenjataan apa pun'. Dalam Laporan Kesamaan yang disebutkan di atas, Yang Mulia Bapak Jānis Kārkliņš mencatat bahwa:

Disarankan agar GGE LAWS mengkatalogkan potensi risiko dan langkah-langkah mitigasi yang harus dipertimbangkan dalam perancangan, pengembangan, pengujian, dan penyebaran sistem persenjataan berbasis teknologi baru di bidang AWS yang mematikan.

Karena Bab 12 buku ini mengeksplorasi unsur kendali manusia yang bermakna, fokus bagian ini adalah pada tiga kewajiban perlindungan negara berdasarkan LOAC, yaitu kewajiban

untuk melakukan tinjauan hukum terhadap senjata baru, kewajiban untuk mengambil tindakan pencegahan dalam serangan, dan kewajiban untuk mengambil tindakan pencegahan terhadap dampak serangan. Hal ini khususnya penting dalam pengerahan senjata yang menunjukkan risiko ketidakpastian, karena senjata-senjata ini dirancang untuk memastikan jaminan prosedural guna menghindari atau meminimalkan risiko.

Tinjauan hukum senjata baru

Negara-negara pihak Protokol Tambahan I, menurut Pasal 36, berkewajiban untuk melakukan tinjauan hukum senjata baru:

Dalam studi, pengembangan, perolehan, atau adopsi senjata, sarana, atau metode peperangan baru, Pihak Peserta Agung berkewajiban untuk menentukan apakah penggunaannya, dalam beberapa atau semua keadaan, akan dilarang oleh Protokol ini atau oleh aturan hukum internasional lain yang berlaku bagi Pihak Peserta Agung tersebut.

ICRC, dalam 'Panduan untuk Tinjauan Hukum Senjata, Sarana, dan Metode Peperangan Baru' tahun 2006, menekankan bahwa kewajiban ini mengalir 'secara logis dari kebenaran bahwa negara dilarang menggunakan senjata, sarana, dan metode peperangan ilegal atau menggunakan senjata, sarana, dan metode peperangan secara ilegal', dan dengan demikian berlaku terlepas dari apakah suatu negara merupakan pihak dalam Protokol Tambahan I. GGE tentang HUKUM memberikan referensi pada kata-kata Pasal. 36 dengan menyatakan bahwa

[s]esuai dengan kewajiban negara-negara berdasarkan hukum internasional, dalam studi, pengembangan, perolehan, atau adopsi senjata, sarana, atau metode peperangan baru, harus dibuat suatu penentuan apakah penggunaannya, dalam beberapa atau semua keadaan, akan dilarang oleh hukum internasional.

Selain itu, laporan yang sama mengakui bahwa negara-negara bebas untuk 'secara independen menentukan cara untuk melakukan tinjauan hukum'. Manfaat berbagi praktik terbaik juga diakui, tidak hanya dalam laporan tersebut, tetapi juga dalam pengajuan masing-masing delegasi. Sebagai langkah ke arah standarisasi praktik negara di bidang tinjauan hukum, Argentina mengajukan 'Kuesioner tentang Mekanisme Tinjauan Hukum Senjata, Sarana, dan Metode Peperangan Baru'. Mengenai hal-hal substantif, kuesioner tersebut menimbulkan pertanyaan penting tentang apakah negara-negara mempertimbangkan pemeriksaan yang telah dilakukan oleh negara lain, apakah mereka melakukan pengujian mereka sendiri, dan apakah mereka pernah menolak perolehan senjata baru.

Menurut Komentar Nasional Australia yang disampaikan kepada GGE tentang LAWS pada tahun 2020, 'memperkuat kepatuhan terhadap HHI yang ada, termasuk melalui tinjauan Pasal 36, adalah cara paling efektif untuk mengelola sistem persenjataan baru, termasuk potensi pengembangan LAWS'. Karena adanya jarak antara masukan manusia langsung pada keputusan penargetan konkret dan tindakan penargetan akhir, tinjauan hukum perlu memeriksa apakah senjata tersebut mampu beroperasi dalam parameter LOAC. Ketika

menyangkut cakupan material tinjauan hukum berdasarkan Pasal 36, ICRC menganggapnya luas dan mencakup 'senjata yang ada yang dimodifikasi sedemikian rupa sehingga mengubah fungsinya, atau senjata yang telah lolos tinjauan hukum tetapi kemudian dimodifikasi'. Untuk AWS yang mengandalkan pengambilan keputusan algoritmik, dapat diharapkan akan ada pembaruan pada parameterinya. Batasan penting adalah 'modifikasi' – yang merupakan pembaruan yang memenuhi syarat sebagai 'modifikasi' senjata dan yang mana yang tidak memenuhi ambang batas tersebut.

Oleh karena itu, tinjauan hukum bukan hanya merupakan pemeriksaan awal, tetapi juga pemeriksaan berkelanjutan terhadap pengembangan dan pengerahan kemampuan militer. Selain menyediakan kerangka kerja yang diperlukan untuk pengenalan senjata baru yang bertanggung jawab, proses ini merupakan sumber informasi tentang sistem, dengan segala risiko dan keterbatasannya. Karena bersifat menghasilkan informasi, proses ini menciptakan kumpulan pengetahuan yang kemudian dapat menjadi masukan bagi pemahaman mereka yang menggunakannya di medan perang. Dengan demikian, kewajiban untuk melakukan tinjauan berdampak pada penilaian kami berdasarkan aturan LOAC yang bersifat *melarang*, karena informasi yang tersedia bagi pembuat keputusan sangat penting dalam menilai keabsahan tindakan mereka.

Tindakan pencegahan dalam serangan

Aspek penting dari perlindungan warga sipil adalah kewajiban untuk mengambil tindakan pencegahan dalam serangan. Aspek dari kewajiban ini berlaku untuk semua operasi militer dan mengharuskan pihak penyerang untuk selalu berhati-hati dalam menyelamatkan penduduk sipil dan objek sipil. Selain itu, kami memiliki daftar kewajiban tindakan pencegahan khusus dalam serangan, termasuk: (a) untuk memverifikasi bahwa target adalah sasaran militer dan bahwa serangan tersebut mematuhi prinsip proporsionalitas; (b) untuk mengambil tindakan pencegahan dalam memilih cara dan metode serangan untuk menghindari, atau meminimalkan, kerugian warga sipil; (c) untuk membatalkan atau menangguhkan serangan jika menjadi jelas bahwa target tersebut adalah sasaran yang dilindungi atau bahwa kerugian insidental akan berlebihan untuk keuntungan militer yang diantisipasi; (d) untuk memberikan peringatan dini yang efektif; dan (e) dalam kasus pilihan antara beberapa sasaran militer yang menghasilkan keuntungan militer serupa, kewajiban untuk memilih sasaran yang diperkirakan menyebabkan kerugian sipil paling sedikit.

Dengan mengambil tindakan pencegahan, pasukan penyerang menciptakan kondisi untuk mematuhi prinsip pembedaan antara warga sipil dan objek sipil, di satu sisi, dan kombatan dan sasaran militer, di sisi lain. Dalam hal ini, mereka merupakan unsur penting dalam arsitektur perlindungan LOAC. Terlepas dari keterkaitan antara kewajiban kehati-hatian dan aturan penargetan lainnya, rezim tindakan pencegahan memiliki signifikansi otonom, karena menetapkan batasannya sendiri terhadap pelaksanaan serangan.

Potensi ketegangan antara kewajiban pihak-pihak yang berkonflik dan pengerahan senjata berbasis AI untuk penargetan langsung muncul melalui kemungkinan pengambilan keputusan senjata secara independen berdasarkan penilaiannya sendiri terhadap perubahan keadaan. Misalnya, jika serangan dilancarkan terhadap sebuah bangunan yang dianggap

digunakan untuk penyimpanan amunisi oleh pihak lawan, informasi baru tentang penggunaan bangunan tersebut dapat diperoleh, yang menunjukkan penggunaan sipilnya, alih-alih militer. Pihak penyerang mungkin memiliki kapasitas untuk membatalkan serangan atau, tergantung pada cara yang digunakan, untuk mengarahkan ulang senjata pasca-pengerahan. Dengan senjata berbasis AI, penilaian informasi baru dan keputusan masing-masing pada akhirnya dapat diserahkan kepada algoritma. Kemungkinan untuk campur tangan dalam pengoperasian senjata juga dapat dibatasi. Oleh karena itu, beberapa negara, termasuk Austria dan Kosta Rika, telah menekankan perlunya manusia mengesampingkan sistem otonom. Di bidang tindakan pencegahan, Amerika Serikat secara konsisten menegaskan bahwa AWS dapat menyediakan serangkaian tindakan baru untuk melindungi warga sipil. Misalnya, menurut Makalah Kerja 2019 yang diserahkan oleh Amerika Serikat kepada GGE, '[m]eskipun risiko korban sipil tidak diharapkan berlebihan dalam kaitannya dengan keuntungan militer yang diharapkan diperoleh, penting untuk mengambil tindakan pencegahan lebih lanjut yang layak. Misalnya, peringatan, pemantauan, dan mekanisme penghancuran diri, penonaktifan diri, atau netralisasi diri adalah tindakan pencegahan yang telah digunakan secara berguna untuk mengurangi risiko jatuhnya korban sipil terkait penggunaan ranjau darat atau laut. Dalam pandangan Amerika Serikat, 'penggunaan sistem otonom bahkan dapat dianggap sebagai tindakan pencegahan tambahan yang harus diambil, konsisten dengan HHI'.

Ada pula pertanyaan apakah negara mampu melaksanakan kewajiban perlindungan mereka bahkan ketika AWS digunakan untuk melakukan serangan. Negara harus mampu menjaga keamanan warga sipil dan objek sipil *secara terus-menerus*, bahkan ketika senjata yang dikerahkan memiliki ruang diskresi sendiri dalam membuat keputusan penargetan. Menerima hal sebaliknya dapat membuat rezim kewajiban kehati-hatian menjadi kurang bermakna. Tentu saja, senjata-senjata ini mungkin saja memiliki kapasitas untuk menyalurkan kewajiban negara: misalnya, sebuah algoritma yang melakukan verifikasi berkelanjutan terhadap informasi baru dan yang juga mampu mengeluarkan peringatan efektif kepada warga sipil. Dalam hal ini, senjata-senjata tersebut mungkin menjadi instrumen yang melaluinya kewajiban untuk menjaga keamanan secara terus-menerus dan kewajiban kehati-hatian spesifik dioperasionalkan. Meskipun pertanyaan apakah teknologi AI akan berkembang sedemikian rupa sehingga memungkinkan hal ini terjadi merupakan pertanyaan empiris, beberapa kesulitan perlu disebutkan. Pertama, meskipun kewajiban kehati-hatian tidak mensyaratkan hasil penghapusan semua kerugian warga sipil, sebagian besar kewajiban tersebut dikondisikan melalui kriteria 'kelayakan', yaitu negara harus 'melakukan *segala sesuatu* yang layak untuk memverifikasi' tujuan mereka dan 'mengambil *semua* tindakan pencegahan yang layak dalam memilih cara dan metode serangan'. Spektrum tindakan yang layak mungkin sulit untuk diprogram sebelumnya. Kedua, perubahan keadaan akan membutuhkan perubahan dalam tindakan pencegahan yang tepat. Misalnya, peringatan dini harus 'efektif', dan efektivitas peringatan tersebut akan bergantung pada berbagai faktor, termasuk pengalaman masa lalu penduduk sipil, pesan yang dikeluarkan oleh pihak lain, kondisi di lapangan, waktu dan kemungkinan untuk memastikan evakuasi yang aman. Ketiga, dalam melaksanakan kewajiban mereka untuk mengambil tindakan pencegahan, pasukan

penyerang dapat mengerahkan senjata yang mempertahankan unsur pengawasan sebagai jaminan terhadap keterlibatan yang tidak diinginkan. Meskipun hal ini mungkin sedikit meredakan kekhawatiran akan penargetan otonom yang tak terkendali, kita tidak boleh melupakan kesulitan yang masih ada: perlunya memastikan pemahaman yang bermakna antara pengawas dan sistem. Yang khususnya penting adalah perlunya memastikan bahwa mereka yang ditugaskan untuk melakukan pengawasan memiliki 'pengetahuan yang diperlukan tentang risiko'. Kebutuhan untuk memahami dinamika manusia-mesin dijelaskan secara eksplisit oleh Inggris dalam Makalah Kerja 2020 yang diserahkan kepada GGE:

mereka yang bertanggung jawab [perlu] memiliki pemahaman yang memadai tentang kemampuan dan keterbatasan sistem persenjataan, serta lingkungan tempat sistem tersebut akan ditempatkan; sebuah poin yang kembali menyoroti pentingnya bentuk interaksi manusia-mesin yang tepat beserta pertimbangan yang lebih luas di seluruh siklus hidup seperti pelatihan personel militer.

Sebagaimana dikemukakan oleh Sassòli dan Quintin, 'kelayakan berkembang melalui pengalaman'. Belajar dari keberhasilan atau kegagalan tindakan pencegahan yang diterapkan sebelumnya adalah kunci pemenuhan kewajiban ini. Hal ini khususnya penting bagi interaksi antara kewajiban untuk mengambil tindakan pencegahan dalam serangan dan kewajiban untuk melakukan tinjauan hukum terhadap senjata. Kewajiban untuk melakukan tinjauan hukum tidak terbatas pada senjata 'baru', tetapi juga berlaku untuk modifikasi teknologi yang sudah ada. Oleh karena itu, perlu ada umpan balik yang bermakna antara pelaksanaan serangan tertentu, pembelajaran dari efektivitas tindakan pencegahan yang diterapkan, kemungkinan perubahan pada algoritma atau fitur model interaksi manusia-mesin, dan tinjauan hukum selanjutnya.

Tindakan pencegahan terhadap dampak serangan

Pasal 58 Protokol Tambahan I memuat kewajiban untuk mengambil tindakan pencegahan terhadap dampak serangan, yang mengharuskan pihak-pihak yang bertikai untuk (1) berupaya memindahkan penduduk sipil, warga sipil perorangan, dan objek sipil di bawah kendali mereka dari sekitar sasaran militer; (2) menghindari penempatan sasaran militer di dalam atau di dekat wilayah yang padat penduduk; dan (3) mengambil tindakan pencegahan lain yang diperlukan untuk melindungi penduduk sipil, warga sipil perorangan, dan objek sipil di bawah kendali mereka dari bahaya yang diakibatkan oleh operasi militer. Ketiga kewajiban tersebut tunduk pada standar 'sejauh yang memungkinkan' yang ditemukan dalam bab ketentuan tersebut. Menurut ICRC, kewajiban untuk mengambil tindakan pencegahan terhadap dampak serangan merupakan bagian dari Hukum Humaniter Internasional kebiasaan. Kewajiban ini dipandang sebagai pelengkap dari kewajiban yang mengikat pasukan penyerang, dan operasi paralel dari kedua kewajiban ini memberikan dasar untuk menerapkan prinsip perbedaan.

Namun, beberapa dari kewajiban ini dapat memberikan beban yang sangat berat pada pihak-pihak yang berkonflik, yang menjelaskan bahasa yang dikondisikan melalui standar kelayakan. Misalnya, kewajiban untuk menghindari menempatkan sasaran militer di dalam

atau di dekat daerah padat penduduk dapat menjadi sangat sulit bagi negara-negara berpenduduk padat atau negara-negara dengan fitur geografis tertentu. Istilah 'kelayakan' tidak didefinisikan dalam Protokol Tambahan I dan harus dinilai tergantung pada keadaan masing-masing kasus. Namun, jelas bahwa kelayakan tidak boleh dinilai setelahnya, dan bahwa keputusan pasukan yang bertahan harus dinilai dengan mengacu pada keadaan dan informasi pada saat itu.

Kesulitan tambahan muncul ketika seseorang mempertimbangkan interaksi antara kewajiban pasukan penyerang dan kewajiban pihak yang bertahan. Kewajiban-kewajiban ini memiliki eksistensi yang independen, dan pelanggaran oleh satu pihak tidak membebaskan pihak lain dari kewajibannya sendiri. Namun, mungkin saja cara satu pihak melaksanakan kewajiban pencegahannya akan memengaruhi efektivitas kewajiban yang diambil oleh pihak lain, atau strategi yang harus diadopsi. Pertanyaan tentang interaksi semacam itu kemungkinan akan mendapatkan perhatian khusus dalam konteks perkembangan teknologi. Dalam artikelnya yang berjudul 'Tindakan Pencegahan terhadap Dampak Serangan di Wilayah Perkotaan', Eric Talbot Jensen secara meyakinkan berpendapat bahwa teknologi baru 'akan memungkinkan pihak pembela memiliki kesadaran situasional yang lebih besar mengenai keberadaan warga sipil, tetapi juga akan meningkatkan kemampuan pihak pembela untuk memisahkan pasukan militer dari warga sipil dan melindungi mereka yang tidak dapat dipisahkan'. Di antara potensi manfaat lain dari pengenalan teknologi tersebut, ia berpendapat bahwa teknologi tersebut dapat digunakan untuk menandai kawasan lindung:

[s]elain itu, penandaan tersebut dapat digunakan untuk menandai kawasan lindung, terutama di medan perang yang dinamis. Jika warga sipil yang terlantar dikumpulkan di sebuah gedung ad hoc, penandaan yang dikirimkan melalui drone dapat digunakan untuk memberi tahu penyerang tentang sifat terlindungi dari gedung tersebut, meskipun hanya sementara.

AI juga dapat digunakan untuk memperingatkan warga sipil melalui 'sinyal visual, seperti lampu kilat atau cat semprot berwarna cerah dengan warna yang telah ditentukan sebelumnya'. Meskipun kemajuan teknologi mungkin benar dapat membantu meminimalkan risiko yang dihadapi warga sipil, risiko baru juga dapat muncul melalui interaksi antara tindakan yang diambil oleh pasukan penyerang dan pasukan bertahan. Sebagai contoh, mari kita pertimbangkan rudal berbasis AI yang dilatih berdasarkan data dari lingkungan tempat rudal tersebut dikerahkan, termasuk kebiasaan warga sipil yang tinggal di sana dan taktik yang digunakan sebelumnya. Jika pihak bertahan memutuskan untuk menggunakan tindakan pemberian sinyal atau penandaan yang baru, yaitu bukan bagian dari informasi yang digunakan untuk melatih algoritma, maka tindakan ini dapat meningkatkan risiko bagi warga sipil, alih-alih mengurangnya. Hal ini karena akan ada faktor lingkungan baru yang dapat bereaksi secara tidak diinginkan. Contoh ini menyoroti perlunya setidaknya pemahaman dasar tentang bagaimana berbagai tindakan dapat memengaruhi fungsi sistem AI. Jika tidak, tindakan pencegahan tertentu tidak hanya dapat menjadi tidak efektif tetapi juga merugikan keselamatan warga sipil.

11.5 KESIMPULAN

Risiko melekat dalam konflik bersenjata. Dampaknya terhadap warga sipil dapat dirasakan dalam berbagai cara, dan kerentanan mereka membutuhkan perlindungan hukum yang kuat. Berdasarkan LOAC, warga sipil dilindungi melalui jaringan kewajiban yang saling terkait dan kompleks. Dengan diperkenalkannya AI, aturan-aturan yang telah mapan ini telah diteliti dengan saksama. Masih banyak pekerjaan yang perlu dilakukan untuk melampaui tahap penegakan umum bahwa hukum internasional berlaku untuk AWS, dan menuju cara-cara konkret penerapannya. Diskusi di GGE tentang LAWS telah berperan penting dalam mempersempit pertanyaan-pertanyaan yang paling menantang dan mengungkap perbedaan pendapat antara interpretasi yang diberikan negara terhadap aturan-aturan hukum internasional tertentu. Meskipun terdapat titik temu pada tingkat umum, titik temu tersebut menjadi kabur ketika seseorang memasuki pertanyaan-pertanyaan konkret tentang penentuan orang-orang yang harus mengendalikan senjata-senjata ini, tahap pengendalian tersebut, potensi kendala atas penggunaannya, dan kekokohan kewajiban perlindungan afirmatif yang mendahului dan menyertai serangan.

Yang tersirat dari analisis perlindungan ini adalah kenyataan bahwa tidak terdapat spesifikasi yang memadai dalam cara pelaksanaan kewajiban perlindungan ini. Meskipun fleksibilitas ini mungkin diperlukan, konkretisasi akan memperkuat perlindungan yang diberikan oleh rezim tersebut. Selain pertanyaan mengenai spesifikasi kewajiban hukum, terdapat tantangan lebih lanjut dalam menilai interaksi antar kewajiban perlindungan, khususnya dampak tindakan yang diambil oleh satu pihak terhadap tindakan yang harus diambil oleh pihak lain.

Semua pertanyaan ini merupakan kunci bagi tanggung jawab negara. Suatu negara bertanggung jawab atas tindakan yang salah secara internasional ketika tindakan yang terdiri dari suatu tindakan atau kelalaian dikaitkan dengan negara tersebut berdasarkan hukum internasional dan merupakan pelanggaran kewajiban internasional negara tersebut. Karena terdapat ketidakpastian mengenai ruang lingkup aturan-aturan tertentu dalam LOAC, pertanyaan apakah suatu negara telah melanggar LOAC menjadi pertanyaan yang sulit untuk dijawab. Menurut Prinsip Panduan (d) dari laporan GGE 2019,

akuntabilitas untuk mengembangkan, menyebarkan, dan menggunakan sistem persenjataan yang baru muncul dalam kerangka CCW harus dipastikan sesuai dengan hukum internasional yang berlaku, termasuk melalui pengoperasian sistem tersebut dalam rantai komando dan kendali manusia yang bertanggung jawab.

Meskipun pernyataan ini menegaskan pentingnya akuntabilitas, rujukan pada 'sesuai dengan hukum internasional yang berlaku' membawa kita kembali pada pertanyaan tentang spesifikasi aturan. Diskusi-diskusi mengenai otonomi ini telah menyoroti area-area dalam LOAC yang belum sepenuhnya ditentukan, khususnya yang berkaitan dengan kesalahan manusia, malfungsi, dan, secara lebih umum, penargetan dalam kondisi berisiko. Terdapat peluang untuk memperjelas aturan-aturan ini. Mengetahui apa yang diharapkan LOAC dari pihak-pihak yang berkonflik merupakan langkah penting untuk memastikan bahwa warga sipil

terlindungi secara efektif dari bahaya operasi militer, bahkan di tengah lanskap konflik yang terus berubah.

BAB 12

IMPLIKASI HUKUM AI DALAM KONFLIK BERSENJATA

12.1 PENDAHULUAN

Sebagaimana dibahas dalam Bab 11, perdebatan mengenai AWS telah mengungkap banyak kelemahan dalam regulasi LOAC dan tanggung jawab negara. AWS juga dapat memengaruhi keseimbangan antara pertimbangan kemanusiaan dan kebutuhan militer. Namun, *mens rea* pelanggaran LOAC dibangun berdasarkan perilaku yang diinginkan pelaku, yang akan sulit dibuktikan dalam konteks AWS. Akibatnya, hukum tentang tanggung jawab negara tidak dapat mengatasi semua tantangan utama yang ditimbulkan oleh AWS. Bab ini menyajikan rezim tanggung jawab negara yang diadopsi oleh NATO dan Amerika Serikat, yang dapat dikatakan meningkatkan efektivitas mekanisme akuntabilitas, transparansi dalam hukum internasional, dan rasa keadilan korban. Model-model ini memungkinkan individu untuk meminta kompensasi atas kerusakan dan kerugian yang terjadi dalam pelaksanaan operasi militer.

Bab ini dibagi menjadi dua bagian utama. Bagian pertama mengkaji dampak hukum pertanggungjawaban negara terhadap aturan-aturan sekunder LOAC, khususnya Pasal-Pasal tentang Tanggung Jawab Negara atas Tindakan-Tindakan yang Melanggar Hukum Internasional tahun 2001 (ARSIWA). Bagian kedua membahas rezim pertanggungjawaban atas dampak permusuhan terhadap pihak sipil, berdasarkan analisis peraturan dari NATO dan Amerika Serikat.

12.2 TANGGUNG JAWAB NEGARA ATAS PELANGGARAN HUKUM PERANG

Tindakan yang Melanggar Hukum Internasional

Menurut Pasal 2 ARSIWA, suatu tindakan yang melanggar hukum internasional terjadi jika dua unsur berikut terpenuhi: (1) tindakan suatu negara melanggar kewajiban internasional negara tersebut, dan (2) tindakan tersebut dapat dikaitkan dengan negara tersebut. Pelanggaran norma primer LOAC dapat menimbulkan pertanggungjawaban atas suatu tindakan yang melanggar hukum internasional. Hal ini karena hukum pertanggungjawaban negara bersifat sekunder terhadap peraturan LOAC, dan berlaku setelah kewajiban primer dilanggar. Namun, pelanggaran LOAC tidak selalu sama dengan kejahatan internasional.

Karakter tanggung jawab negara – baik perdata maupun pidana – tidak sepenuhnya jelas. Misalnya, dalam Laporan Kelima tentang Tanggung Jawab Negara, Arangio-Ruiz menyimpulkan bahwa tanggung jawab internasional harus dilihat dari perspektif perdata dan pidana. Aspek perdata dari tanggung jawab berarti bahwa dalam beberapa kasus dapat mengarah pada kompensasi, sementara aspek pidana dapat mengarah pada penuntutan. Namun, penting untuk dicatat bahwa, pada prinsipnya, negara tidak berada di bawah yurisdiksi wajib pengadilan internasional. Sekalipun yurisdiksi tersebut berlaku, pengadilan internasional tidak menyediakan mekanisme penegakan hukum. Meskipun terdapat mekanisme sanksi, mekanisme tersebut didasarkan pada Piagam PBB, dan berorientasi pada perdamaian dan

keamanan internasional. Munculnya rezim baru tanggung jawab internasional (yaitu tanggung jawab internasional dan tanggung jawab pidana individu) telah memberikan pandangan baru tentang tanggung jawab negara, dengan dua yang pertama dianggap sebagai fungsi dari yang terakhir. Meskipun demikian, rezim-rezim tersebut telah meningkatkan saling ketergantungan dan munculnya isu-isu global yang sangat memengaruhi sifat tanggung jawab negara. Munculnya teknologi-teknologi baru, dan AWS di antaranya, menantang nilai-nilai universal umat manusia, sehingga menuntut respons global.

Sifat tanggung jawab negara selanjutnya bergantung pada kewajiban yang dilanggar. Berdasarkan Pasal 42 ARSIWA, kewajiban internasional dibagi ke dalam kategori-kategori berikut: (1) yang dibebankan kepada negara yang dirugikan secara individual; (2) yang dibebankan kepada sekelompok negara (*obligations erga omnes partes*); dan (3) yang dibebankan kepada komunitas internasional secara keseluruhan (*obligations erga omnes*). Sebuah kategori independen dari kewajiban internasional muncul berdasarkan Pasal 40 ARSIWA, yang mengacu pada norma peremptory hukum internasional umum dan pelanggaran serius yang terkait terhadap kewajiban tersebut. Oleh karena itu, karakter pelanggaran menentukan subjek dan ruang lingkup reaksi yang dapat diterima terhadap pelanggaran tersebut. Dalam kasus-kasus di mana kewajiban itu dimiliki oleh negara yang terluka secara individual, hanya negara yang terluka yang memiliki kepentingan hukum dalam realisasi tanggung jawab. Jika ada pelanggaran kewajiban yang dimiliki oleh komunitas internasional secara keseluruhan, negara-negara lain (bukan hanya yang terluka) dapat meminta tanggung jawab. Akibatnya, disarankan bahwa penggunaan teknologi tertentu dalam konflik bersenjata, seperti 'sistem senjata yang memilih dan menyerang target berdasarkan pemrosesan sensor', adalah masalah yang menjadi perhatian komunitas internasional secara keseluruhan. Contoh senjata tersebut adalah: sistem Harpy buatan Israel (dijelaskan sebagai AWS 'dirancang untuk mendeteksi, menyerang, dan menghancurkan pemancar radar'), Sistem Senjata Dekat Phalanx AS untuk kapal penjelajah kelas Aegis, pesawat tanpa awak tempur bertenaga jet Taranis Inggris, atau Samsung Techwin. Daftar kewajiban LOAC yang merupakan kewajiban yang dimiliki oleh komunitas internasional dapat ditemukan dalam pasal. Pasal 85(1) Protokol Tambahan I, yang menjelaskan pelanggaran berat Konvensi Jenewa I–IV. Pelanggaran tersebut mencakup, misalnya, pembunuhan yang disengaja terhadap orang-orang yang dilindungi, penyebab yang disengaja atas penderitaan yang hebat, atau cedera serius pada tubuh atau kesehatan, penghancuran dan perampasan properti yang luas dan tidak dapat dibenarkan. Berdasarkan laporan ILC tahun 2006, sebagian besar LOAC merupakan kewajiban *erga omnes* karena sifatnya yang non-timbal balik dan perlindungan yang dilakukan demi kepentingan semua negara. Namun, ILC tidak memberikan contoh kewajiban semacam itu di bidang LOAC. Dalam beberapa kasus, ICJ merujuk pada 'beberapa' atau 'banyak sekali' aturan LOAC yang menciptakan kewajiban *erga omnes*. Tidak satu pun dari aturan tersebut yang memberikan daftar lengkap mengenai pelanggaran aturan LOAC mana yang menyebabkan terciptanya kewajiban semacam ini. Situasi ini menimbulkan ambiguitas dan ketidakjelasan yang semakin meningkat atas kewajiban yang dimiliki oleh komunitas internasional secara keseluruhan.

Meskipun demikian, kewajiban LOAC tersebut berasal dari Pasal 1 I–IV Konvensi Jenewa, yang menetapkan kewajiban untuk memastikan penghormatan terhadap LOAC. Kewajiban ini awalnya dianggap sebagai tugas yang ditujukan kepada angkatan bersenjata dan otoritas publik suatu negara. Saat ini, kewajiban *erga omnes* dianggap membebaskan kewajiban untuk menuntut pelaksanaan oleh negara lain. Dalam beberapa kasus, kewajiban ini mencakup kewajiban untuk memastikan penghormatan dari negara lain. Dalam kasus AWS, kewajiban untuk memastikan penghormatan dari negara lain mencakup langkah-langkah yang sah untuk tidak mentransfer AWS ke negara yang melanggar LOAC, atau pengenaan sanksi ekonomi terhadap negara yang melanggar. Oleh karena itu, suatu negara berkewajiban untuk tidak mendorong orang atau kelompok untuk bertindak melanggar LOAC.

Kewajiban Negara Terkait Penggunaan Aws Dalam Konflik Bersenjata

Kewajiban negara dalam LOAC dapat dibagi menjadi dua kategori. Pertama, ada kewajiban positif untuk menghormati LOAC, yang timbul terlepas dari keberadaan konflik bersenjata. Kedua, ada kewajiban negatif yang harus dipenuhi ketika konflik bersenjata terjadi. Pelanggaran kategori kedua dapat menyebabkan pelanggaran berat terhadap LOAC.

Sesuai dengan pasal 1 Konvensi Jenewa I–IV, serta pasal 1(1) Protokol Tambahan I, negara berkewajiban untuk menghormati dan memastikan penghormatan terhadap perjanjian-perjanjian ini dalam segala keadaan. Kewajiban ini tidak hanya berasal dari perjanjian, tetapi juga dari hukum kebiasaan internasional. Suatu negara harus memastikan kepatuhan terhadap LOAC dengan mencegah pelanggaran, dan dengan menuntut dan menghukum pelanggaran LOAC yang dilakukan oleh warga negaranya, atau yang dilakukan di wilayah negara tersebut. Dalam konteks AWS, jaminan kepatuhan dapat dicapai dengan melakukan peninjauan senjata berdasarkan pasal. 36 Protokol Tambahan I, yang menyatakan bahwa:

dalam studi, pengembangan, perolehan, atau adopsi senjata, sarana, atau metode peperangan baru, Pihak Kontrak Tinggi berkewajiban untuk menentukan apakah penggunaannya, dalam beberapa atau semua keadaan, akan dilarang oleh Protokol ini atau oleh aturan hukum internasional lain yang berlaku bagi Pihak Kontrak Tinggi tersebut.

Meskipun kewajiban ini merupakan perlindungan yang diperlukan untuk kepatuhan LOAC sehubungan dengan pengembangan dan penggunaan AWS, tidak ada mekanisme kelembagaan yang memungkinkan verifikasi apakah suatu negara memenuhi kewajiban tersebut (kecuali untuk pasal 84 Protokol Tambahan I tentang kewajiban untuk berkomunikasi). Selain itu, pasal 36 Protokol tidak mencegah entitas swasta mengembangkan teknologi persenjataan karena kewajibannya hanya dibebankan kepada negara. Meskipun demikian, dalam diskusi tentang legalitas AWS, banyak negara mengajukan prosedur peninjauan yang berkaitan dengan sarana dan metode peperangan baru. AWS memang memiliki referensi langsung dalam undang-undang AS. Pedoman Peninjauan Sistem Senjata Otonom atau Semi-Otonom Tertentu memberikan kewajiban ganda untuk menilai legalitas AWS. Pemeriksaan ulang diperlukan baik dalam fase pengembangan maupun penerapan. Oleh

karena itu, negara-negara tidak bersedia menyajikan hasil prosedur peninjauan mereka terkait AWS; namun, prosedurnya sendiri memang dikomunikasikan secara publik.

Negara-negara berkewajiban untuk memperkenalkan jaminan prosedural yang berlaku untuk pelanggaran LOAC. Jaminan ini mencakup: (1) sanksi pidana yang efektif; (2) penuntutan atau ekstradisi tersangka pelaku; dan (3) yurisdiksi universal untuk pelanggaran LOAC yang serius. Sehubungan dengan kewajiban pertama, sanksi pidana yang efektif harus diperkenalkan ke dalam undang-undang domestik. Jika negara mengembangkan atau memiliki AWS, yang rencananya akan digunakan dalam konflik bersenjata, negara tersebut harus mempertimbangkan untuk mengesahkan undang-undang yang relevan yang akan mengatur penggunaan teknologi tersebut. Dalam konteks AWS, negara-negara yang mengembangkan dan menerapkan teknologi tersebut dalam operasi militer termasuk Israel dan Amerika Serikat.

Poin kedua mengacu pada aspek penuntutan dan ekstradisi yang memaksakan kewajiban *erga omnes partes*, yaitu ditetapkan untuk melindungi kepentingan kolektif. Contoh bagus dari hal ini adalah kasus Israel, yang didakwa dengan kejahatan merancang, memproduksi, dan menjual AWS ke negara yang tidak dikenal. Meskipun demikian, ekstradisi merupakan aspek yang menantang karena beberapa pengecualian yang dapat diterapkan, serta perlunya perjanjian ekstradisi. Oleh karena itu, disarankan bahwa kewajiban untuk menuntut pelanggaran harus diutamakan daripada ekstradisi. Masalah-masalah ini bahkan lebih dalam dalam kasus-kasus dugaan kejahatan perang dengan AWS yang terlibat dalam tindakan yang merupakan pelanggaran LOAC. Aspek ini telah dibahas dalam bab sebelumnya. Praktik-praktik yang berkaitan dengan senjata lain mengungkapkan ketidakefisienan perlindungan internasional terhadap penggunaan sistem persenjataan secara ilegal. Misalnya, meskipun ada serangan dengan senjata pembakar dan kimia di Suriah, tidak ada informasi tentang upaya untuk meminta pertanggungjawaban negara atau mengadili para pelaku yang diduga. Hal ini membuat para korban tidak memiliki alat untuk melawan badan-badan negara yang bertanggung jawab untuk menentukan perlunya penyelidikan suatu kasus.

Poin terakhir, mengenai aturan yurisdiksi universal, mencakup pelaksanaan yurisdiksi untuk pelanggaran LOAC yang serius di luar batas negara,³⁴ terlepas dari hubungan antara negara dan lokasi, pelaku, atau tindakan. Yurisdiksi ini ditegakkan demi kepentingan bersama semua negara. Yurisdiksi universal tidak hanya telah dimasukkan ke dalam Konvensi Jenewa I-IV (masing-masing pasal 49, 50, 129, 146), tetapi juga ke dalam perjanjian lain yang berlaku untuk situasi konflik bersenjata. Oleh karena itu, yurisdiksi universal juga mencakup tindakan yang dilakukan dengan menggunakan AWS dalam konflik bersenjata. Sebagian dari tindakan tersebut akan melibatkan pelanggaran berat atau pelanggaran LOAC yang serius, dengan asumsi bahwa tindakan tersebut dilakukan dengan sengaja dan ditujukan terhadap penduduk atau objek sipil. Rasionalnya adalah bahwa penggunaan AWS dianggap sebagai penggunaan sarana dan metode peperangan. Pelaksanaan yurisdiksi universal berkontribusi pada budaya tanggung jawab dengan mencegah penghindaran tanggung jawab melalui inefisiensi prosedural dalam tatanan hukum nasional. Ini merupakan alat pencegahan terhadap kecelakaan dengan AWS yang dilakukan oleh personel yang tidak terlatih secara memadai atau di lingkungan di mana sistem itu sendiri belum teruji secara memadai.

Undang-undang tentang tanggung jawab negara tidak memberikan informasi yang tepat mengenai pelanggaran LOAC mana yang merupakan tanggung jawab negara dan konsekuensi apa yang harus diterapkan. Meskipun tanggung jawab negara atas penggunaan AWS awalnya diangkat dalam diskusi, argumen tersebut tampaknya digunakan sebagai tanggapan sementara untuk mengatasi pelanggaran LOAC. Namun, hal itu bukanlah satu-satunya jawaban yang mungkin atas dampak perilaku negara dalam konflik bersenjata. Klarifikasi tertentu dapat ditemukan dalam paragraf 3 Prinsip dan Pedoman Dasar 2005, yang menekankan bahwa negara memiliki kewajiban untuk 'memberikan pemulihan yang efektif kepada korban, termasuk reparasi', yang harus memadai, efektif, dan cepat. Namun demikian, instrumen hukum ini hanya membahas kategori pelanggaran LOAC serius 'yang, berdasarkan sifatnya yang sangat berat, merupakan penghinaan terhadap martabat manusia'. Dengan kata lain, instrumen ini tidak mencakup pelanggaran LOAC lainnya ketika suatu tindakan oleh AWS tidak termasuk dalam kategori ini. Tindakan-tindakan ini mengasumsikan perilaku manusia yang disengaja dan langsung, sehingga mengecualikan kesalahan atau kekeliruan AI yang tidak disengaja atau tidak terduga. Tidak ada satu jawaban tunggal karena ada banyak jenis AWS yang berbeda, oleh karena itu, setiap kasus memerlukan pemeriksaan individual.

Elemen lain dari tindakan yang salah secara internasional adalah atribusi tindakan tersebut kepada negara. Tindakan tersebut harus dilakukan oleh individu (agen atau perwakilan) yang hubungannya dengan negara memungkinkan tindakan mereka dikaitkan dengan negara itu. Oleh karena itu, pertanyaan mendasar untuk aturan atribusi adalah pembentukan hubungan kelembagaan atau faktual antara suatu negara dan orang atau entitas individu yang melakukan tindakan tersebut. Suatu negara memikul tanggung jawab atas perilaku tidak hanya angkatan bersenjatanya, tetapi juga organ-organ negara lainnya dan, dalam hal pengecualian, individu-individu swasta. Pasal 4–11 ARSIWA memberikan daftar entitas, yang perilakunya dihitung untuk perilaku suatu negara, yaitu: (1) organ-organ suatu negara; (2) orang atau entitas yang menjalankan unsur-unsur otoritas pemerintahan; (3) organ-organ yang ditempatkan di bawah kendali suatu negara oleh negara lain; (4) orang atau sekelompok orang yang diarahkan atau dikendalikan oleh suatu negara; (5) gerakan pemberontakan atau gerakan lainnya; dan (6) orang atau entitas yang tindakannya diakui oleh suatu negara sebagai tindakannya sendiri. Dalam kasus AWS, terkadang sulit untuk menetapkan hubungan dengan negara.

Sehubungan dengan angkatan bersenjata, berdasarkan LOAC, suatu negara bertanggung jawab atas perilaku anggota angkatan bersenjatanya. Hal ini khususnya penting ketika tindakan mereka dilakukan bertentangan dengan perintah atau instruksi. Selain itu, negara juga bertanggung jawab atas perilaku orang atau entitas yang bertindak di bawah instruksi, arahan, atau kendali negara tersebut. Aturan dalam konflik bersenjata ini merujuk, misalnya, kepada kelompok bersenjata yang didukung oleh satu negara dan berperang melawan negara lain. Uji hukum untuk atribusi dapat ditemukan dalam *kasus Nikaragua*, di mana ICJ menyatakan bahwa harus dibuktikan bahwa negara secara efektif mengendalikan operasi militer dan paramiliter suatu kelompok bersenjata. Uji tersebut tidak terpenuhi jika negara hanya membiayai, mengorganisir, melatih, mempersenjatai, mengirimkan peralatan,

merencanakan operasi, dan memilih target. Akibatnya, pengiriman AWS ke kelompok bersenjata atau pilihan target yang telah diprogram sebelumnya ke dalam AWS oleh negara pengirim tidak akan cukup untuk memenuhi uji atribusi. Negara harus secara efektif menginstruksikan dan mengarahkan atau mengendalikan kelompok bersenjata yang melakukan pelanggaran LOAC.

Cakupan kendali negara pengirim atas AWS akan berkontribusi pada tanggung jawabnya, jika kendali ini sebenarnya efektif. Kategori instruksi yang diberikan kepada entitas swasta atau individu yang mengembangkan atau mengendalikan AWS sangat menarik. Negara yang membeli AWS dari individu swasta akan memberikan spesifikasi AWS, yang kemudian dapat dinilai dari perspektif tanggung jawab. Pada tahun 2018, kontraktor pertahanan Israel, Aeronautics Ltd menunjukkan kemampuan pesawat tanpa awak bunuh diri Orbiter 1K mereka dengan menyerang angkatan bersenjata Armenia atas nama Azerbaijan. Israel menyatakan akan menuntut para eksekutif perusahaan dan beberapa karyawan lain yang terlibat. Selain itu, Israel menangguhkan lisensi ekspor perusahaan tersebut. Uji dari kasus *Nicaragua* kemudian dikembangkan oleh ICJ dalam kasus *Genosida*, di mana pengadilan menyatakan bahwa harus ada hubungan antara tindakan organ suatu negara atau kelompok bersenjata yang secara faktual bertindak sebagai organ negara tersebut dan tanggung jawab negara.

Akibatnya, tindakan seorang individu dapat dikaitkan dengan negara jika individu tersebut, pada kenyataannya, bertindak dalam ketergantungan penuh pada negara. Oleh karena itu, ada kesenjangan dalam ketentuan yang berkaitan dengan tanggung jawab negara dalam hal penyerahan tunggal AWS dan kurangnya kontrol efektif terhadap mereka yang dilengkapi dengan AWS oleh negara. Batasan dalam kendali non-efektif dan efektif atas AWS yang diserahkan kepada kelompok bersenjata akan sulit dibuktikan. Uji kendali yang diadopsi oleh ICJ mempersempit tanggung jawab negara secara eksklusif pada perilaku spesifik individu yang kepadanya negara menjalankan kendali efektif. Ini berarti bahwa setiap perilaku memerlukan pemeriksaan terpisah.

Oleh karena itu, penggunaan AWS oleh kelompok bersenjata non-negara akan dinilai dari perspektif kendali efektif. Tantangan lain terkait dengan perusahaan militer dan keamanan swasta yang akan mengirim atau melatih angkatan bersenjata tentang cara menggunakan AWS. Argumen untuk menghubungkan perilaku individu yang diinstruksikan oleh negara untuk membuat atau mengendalikan AWS adalah bahwa negara tidak dapat menghindari tanggung jawab dengan benar-benar mendelegasikan fungsi mereka kepada individu swasta. Jika kontraktor secara fungsional terkait dengan suatu negara, negara ini memiliki kewajiban untuk memastikan penghormatan terhadap LOAC oleh kontraktor. Akhirnya, sesuai dengan pasal 16 ARSIWA, suatu negara bertanggung jawab untuk membantu atau mendukung perilaku melanggar hukum negara lain. Atribusi dalam kasus ini muncul jika negara mengetahui tindakan tersebut dan bantuan atau asistensi tersebut dimaksudkan untuk memfasilitasi tindakan tersebut. Ketentuan ini dapat sangat berharga bagi transfer AWS ke negara lain atau bagi pembagian pengawasan terkait AI. Daftar Senjata Konvensional PBB menyajikan transfer teknologi dalam sistem persenjataan antarnegara, yang juga mencakup kategori pesawat tempur dan kendaraan udara nirawak. Menurut Daftar PBB, daftar ini mencakup lebih dari 90%

perdagangan senjata global. Meskipun demikian, tidak ada kategori yang secara khusus merujuk pada teknologi senjata AI.

12.3 REZIM PERTANGGUNGJAWABAN

Alasan Pertanggungjawaban

Pertanggungjawaban memberikan kewajiban untuk membayar kompensasi atas kerugian yang disebabkan oleh tindakan yang tidak dilarang dalam hukum internasional. Kewajiban ini ditentukan oleh: (1) adopsi perjanjian internasional yang memperkenalkan kewajiban kompensasi; (2) munculnya aturan kebiasaan tentang kewajiban untuk memberikan kompensasi; dan (3) penerimaan perkembangan progresif aturan hukum internasional yang menetapkan risiko pertanggungjawaban. Tidak ada perjanjian yang secara langsung mengatur pertanggungjawaban atas dampak penggunaan sistem persenjataan apa pun; kasus-kasus semacam itu diatur di dalam negeri. Sifat pertanggungjawaban yang lazim juga kontroversial. Ada penulis yang mengusulkan pertanggungjawaban sebagai aturan kebiasaan yang baru muncul, tetapi hanya dalam kaitannya dengan rezim tertentu, misalnya, untuk kerugian lintas batas.

Karena tidak ada peraturan internasional mengenai reparasi dan kompensasi atas dampak permusuhan yang tidak dilarang dalam hukum internasional, negara-negara sendiri memberlakukan undang-undang yang relevan mengenai gugatan perang. Yang terakhir tidak diakibatkan oleh tindakan terlarang, tetapi berorientasi pada kerugian atau kerusakan. NATO telah memberikan latar belakang kelembagaan bagi para korban dampak operasi militer yang dipimpin NATO. Amerika Serikat juga mengadopsi undang-undang tentang kompensasi atas dampak operasi militer. Sementara kedua peraturan tersebut mencakup kerugian atau kerusakan yang tidak secara langsung diakibatkan oleh perilaku permusuhan, keduanya menangani tindakan kelalaian. Hal ini dapat memiliki dampak potensial pada penggunaan AWS, jika dampaknya tidak diprediksi dan tidak terbatas pada target militer. Karena karakter konflik bersenjata kontemporer yang berubah dan dinamis, terkadang sulit untuk membedakan antara dampak langsung permusuhan dan dampak operasi militer lainnya.

Model NATO

Meskipun NATO telah diberikan kepribadian hukum internasional, tanggung jawab atau liabilitas organisasi internasional masih samar. Dalam konteks AWS, dapat berupa NATO atau negara yang menyumbangkan pasukan militernya. Untuk menyelesaikan masalah ini, pasal. Pasal 7 dari Rancangan Pasal tentang Tanggung Jawab Organisasi Internasional mengusulkan uji 'kendali efektif' yang harus dibuktikan bahwa organisasi tersebut menjalankan kendali efektif atas perilaku suatu organ atau agen yang ditempatkan di bawah kendali organisasi tersebut. Hal ini dapat diterapkan pada penggunaan AWS dalam operasi yang dipimpin NATO jika, misalnya, suatu negara menyerahkan AWS yang kemudian digunakan dan dikendalikan oleh pasukan yang dipimpin NATO.

Sebuah studi kasus keterlibatan NATO di Afghanistan menggambarkan masalah pertanggungjawaban dengan baik. Untuk menghindari kaitan dengan pertanggungjawaban, istilah 'kompensasi' sering kali diganti dengan '*ex gratia*', ganti rugi pertempuran, pembayaran

kehormatan atau solatium. Akibatnya, Pasukan Bantuan Keamanan Internasional (ISAF) di Afghanistan membentuk apa yang disebut Sel Pelacakan Korban Sipil. Tujuannya adalah untuk mengumpulkan dan memelihara data tentang kerugian yang ditimbulkan pada warga sipil (tetapi bukan pada properti mereka). Karena kurangnya transparansi dalam prosedur militer yang mengarah pada identifikasi unit yang bertanggung jawab, warga sipil biasanya tidak menyadari peran Sel. Karena ketidakefektifannya, Sel kemudian digantikan oleh Tim Mitigasi Korban Sipil. Tim ini bertanggung jawab atas koordinasi studi dan rekomendasi khusus subjek kepada rantai komando, serta koordinasi kelompok kerja yang membahas penetapan pedoman dan prosedur operasi standar terkait proses pelacakan kerugian warga sipil. Kelompok kerja tersebut terdiri dari anggota ISAF dan organisasi internasional serta Afghanistan. Selain itu, negara-negara NATO mengadopsi seperangkat Pedoman yang tidak mengikat tentang Pembayaran Moneter kepada Korban Sipil di Afghanistan. Berdasarkan Pedoman tersebut, ketika menentukan respons yang tepat terhadap korban sipil tertentu, angkatan bersenjata harus mempertimbangkan kebiasaan dan norma setempat, termasuk potensi pembayaran *ex gratia*. Pedoman tersebut bermanfaat di Afghanistan, tetapi tidak diadopsi dalam konflik-konflik selanjutnya, misalnya di Libya.

Praktik AS Dalam Menangani Dampak Permusuhan

Peraturan AS dapat ditemukan dalam Undang-Undang Klaim Asing (FCA), yang awalnya diadopsi untuk meningkatkan hubungan antara angkatan bersenjata AS dan warga sipil di luar negeri. Undang-undang ini memperkenalkan klaim kompensasi hingga Rp. 500.000.000. Seseorang dapat mengklaim kompensasi atas cedera pribadi, kematian, atau kerusakan properti yang disebabkan oleh aktivitas angkatan bersenjata AS yang tidak terkait dengan pertempuran. FCA memberi wewenang kepada Menteri Pertahanan untuk menunjuk Komisi Klaim Asing *ad hoc* untuk menangani kasus-kasus tersebut. Komisi menerapkan hukum dan kebiasaan setempat di negara tempat klaim muncul atau, jika penggugat tinggal di tempat lain, negara tempat tinggal. Putusan Komisi bersifat final.

Meskipun sistem ini memberikan kemungkinan kompensasi, sistem ini memiliki beberapa kelemahan. Pertama, sistem ini bersifat *ad hoc*, yang berarti Komisi beroperasi secara non-permanen. Kedua, berdasarkan § 2734(b)(3) FCA, model kompensasi AS menyediakan apa yang disebut 'pengecualian pertempuran', yang berarti kompensasi tidak dapat diberikan untuk kerugian atau kerusakan yang diakibatkan secara langsung atau tidak langsung dari situasi pertempuran atau dari persiapan langsung untuk pertempuran. Definisi 'pertempuran' dalam praktiknya juga merupakan masalah yang menantang, bahkan untuk Komisi Klaim Asing. Namun, perlu dicatat bahwa banyak klaim berbasis FCA di Irak menangani kerusakan yang disebabkan oleh pelepasan senjata secara lalai, yang penting dalam konteks AWS.

Selain FCA, ada sumber kompensasi alternatif, yaitu sistem *ad hoc* pembayaran *solatia* dan pembayaran belasungkawa di bawah Program Tanggap Darurat Komandan. Ini adalah pembayaran *solatia* berdasarkan jumlah yang telah ditentukan sebelumnya, untuk kematian, cedera, atau kerusakan properti yang disebabkan oleh angkatan bersenjata AS selama pertempuran. Seorang individu (korban langsung atau keluarga korban) dapat mengajukan

permohonan. Pembayaran ditentukan sesuai dengan adat istiadat setempat. Lebih lanjut, pembayaran solatia dan klaim berbasis FCA saling mengecualikan. Sebagaimana ditunjukkan dalam Buku *Panduan Hukum Operasional tahun 2018*, 'individu atau unit yang terlibat dalam kerusakan tidak memiliki kewajiban hukum untuk membayar'. Jenis kompensasi ini hanyalah ungkapan niat baik. Dana Program Tanggap Darurat Komandan telah digunakan untuk melakukan pembayaran belasungkawa di Afghanistan dan Irak sebagai pengakuan atas kerugian. Dampak tak terduga dari penggunaan AWS dalam permusuhan yang tidak termasuk pelanggaran LOAC akan ditangani dalam rezim ini. Istilah 'tak terduga' merujuk, misalnya, pada situasi ketika AWS membuat keputusan secara otonom atau ketika keadaan yang menjadi dasar keputusan AWS berubah.

12.4 KESIMPULAN

Penegakan LOAC terutama difokuskan pada tanggung jawab pidana individu. Namun, para korban, yang seharusnya dilindungi oleh LOAC, hanya memiliki sedikit langkah efektif untuk mengatasi kerugian atau kerusakan yang disebabkan oleh penggunaan AWS dalam konflik bersenjata. Meskipun LOAC memberikan tanggung jawab negara melalui kewajiban untuk membayar kompensasi atas beberapa pelanggaran LOAC, tanggung jawab tersebut akan sulit diterapkan dalam menangani penggunaan AWS, terutama karena masalah atribusi tindakan AWS kepada negara.

Perilaku otonom AWS dapat tidak terduga dan sangat berbahaya pada saat yang bersamaan. Penerimaan risiko ini oleh negara harus diikuti dengan penerimaan tanggung jawab atas beberapa tindakan yang tidak dilarang oleh hukum internasional. Program kompensasi untuk dampak permusuhan berkontribusi pada keberhasilan umum, dan dukungan sipil untuk, operasi militer. Di tingkat domestik, program-program tersebut, meskipun tidak secara tegas membahas penggunaan AWS, telah ada selama lebih dari 60 tahun. Hal ini menyebabkan sistem perlindungan yang berbeda, dan terkadang berbeda-beda, bagi korban, karena bergantung pada karakter konflik dan pihak-pihak yang terlibat. Lebih lanjut, program kompensasi menghadapi tantangan signifikan dalam menerapkan hukum nasional, termasuk interpretasi pengecualian pertempuran atau zona pertempuran. Semua ini semakin membatasi penerapannya pada kasus AWS.

BAB 13

KENDALI MANUSIA ATAS SENJATA OTONOM

13.1 PENDAHULUAN

Persyaratan kendali manusia atas penggunaan sistem senjata otonom (AWS) merupakan topik yang banyak dibahas, namun masih sulit dipahami. Banyak advokat kebijakan dan akademisi berpendapat bahwa AWS berpotensi mematuhi hukum konflik bersenjata (LOAC) hanya jika dipandu oleh kendali manusia, dan hal ini telah dibahas dalam bab-bab sebelumnya. Mereka berpendapat bahwa persyaratan kendali manusia harus setara dengan prinsip-prinsip kunci LOAC lainnya, yaitu prinsip pembedaan, proporsionalitas, kemanusiaan, dan kebutuhan militer. Pemerintah AS telah mengadopsi Direktif 3000.09 tentang AWS, yang menyatakan bahwa sistem senjata otonom 'harus dirancang untuk memungkinkan komandan dan operator *menerapkan tingkat pertimbangan manusia* yang sesuai dalam penggunaan kekuatan'. Rumusan AS tentang peran manusia dalam penggunaan AWS telah mendapatkan dukungan yang signifikan, dan para penulis mulai menyamakan konsep ini dengan persyaratan 'kendali manusia'. Namun, pemerintah AS tidak mendukung deskripsi tersebut dan menentang upaya internasional apa pun untuk mengatur AWS berdasarkan konsep 'kendali manusia'. Bab ini mengeksplorasi lebih dalam perbedaan antara konsep 'penilaian manusia' dan 'kendali manusia' dengan mempelajari asumsi-asumsi yang mendasari konsep-konsep kebijakan yang tampaknya serupa ini. Bab ini juga mengidentifikasi poin-poin penting yang telah diabaikan dalam problematisasi kebijakan Departemen Pertahanan AS (DoD) tetapi tetap merupakan pertimbangan penting yang dapat menginformasikan penilaian kritis terhadap pendekatan DoD AS.

13.2 MASALAH DALAM KEBIJAKAN PERTAHANAN AS

Kebijakan AS yang tercermin dalam Direktif 3000.09 menguraikan tiga jenis sistem robotik yang mendapatkan 'lampu hijau' untuk disetujui dalam kebijakan tersebut. Ketiga jenis tersebut adalah: (1) senjata semi-otonom, seperti amunisi homing; (2) senjata otonom defensif yang diawasi, seperti sistem senjata Aegis berbasis kapal; dan (3) senjata otonom non-mematikan dan non-kinetik, seperti peperangan elektronik untuk mengacaukan radar musuh. Ketiga kelas AWS ini digunakan secara luas saat ini, dan kebijakan tersebut menegaskan bahwa pengembang dapat menggunakan otonomi sesuai dengan praktik yang ada. Senjata-senjata ini tunduk pada aturan akuisisi normal dan tidak memerlukan persetujuan tambahan. Namun, senjata masa depan apa pun yang akan menggunakan otonomi dengan cara baru di luar ketiga jenis tersebut mendapatkan 'lampu kuning'. Sistem-sistem tersebut tunduk pada proses peninjauan yang panjang, dengan fokus utama pada pengujian dan evaluasi, baik sebelum pengembangan maupun pengerahan.

Cara baru yang potensial dalam menggunakan otonomi yang berada di luar contoh yang ditentukan mengacu pada segala jenis senjata otonom yang digunakan untuk *tujuan ofensif*. Sementara AWS ofensif tidak dilarang oleh Direktif 3000.09, pembatasan tambahan

bertujuan untuk mengurangi risiko yang terkait dengan potensi pengembangan senjata semacam itu. Alasan diberlakukannya pembatasan ini adalah untuk 'meminimalkan probabilitas dan konsekuensi kegagalan dalam sistem persenjataan otonom dan semi-otonom yang dapat menyebabkan *keterlibatan yang tidak disengaja*'. Keterlibatan yang tidak disengaja didefinisikan sebagai 'penggunaan kekuatan yang mengakibatkan kerusakan pada orang atau objek yang tidak dimaksudkan oleh operator manusia untuk menjadi target operasi militer AS'. Akibatnya dapat menyebabkan 'tingkat kerusakan kolateral yang tidak dapat diterima' yang bertentangan dengan hukum perang atau Aturan Keterlibatan (ROE). Delegasi AS untuk PBB memberikan contoh serangan tidak disengaja yang menewaskan warga sipil, atau pasukan sekutu akan dianggap sebagai 'keterlibatan yang tidak disengaja' berdasarkan Direktif DoD 3000.09.

Kegagalan didefinisikan sebagai 'degradasi atau hilangnya fungsionalitas yang dimaksudkan atau ketidakmampuan sistem untuk berfungsi sebagaimana mestinya atau dirancang'. Direktif 3000.09 menyatakan bahwa kegagalan dapat disebabkan oleh berbagai penyebab, misalnya kegagalan interaksi manusia-mesin atau serangan siber. Dewan Sains Pertahanan (DSB) menetapkan bahwa tantangan utama untuk adopsi teknologi tersebut secara lebih luas adalah masalah kepercayaan. Laporan tersebut menyatakan bahwa meskipun produsen senjata menerapkan atribut kepercayaan seperti tingkat kompetensi, keandalan, dan integritas yang tinggi pada tingkat desain, masih terdapat masalah dengan kepercayaan operasional yang terkait dengan penggunaan senjata tersebut. Tantangan utamanya adalah AWS mungkin memiliki sensor dan sumber data yang berbeda dari rekan satu tim manusianya, sehingga mereka mungkin beroperasi dengan asumsi kontekstual yang berbeda dari lingkungan operasional. Lebih lanjut, karena sistem tersebut bersifat belajar mandiri, kapabilitasnya dapat bervariasi di berbagai lingkungan, dan 'proses penalaran' mereka mungkin mengambil jalur yang berbeda dari pengambil keputusan manusia.

Namun, terlepas dari bahaya-bahaya ini, Direktif 3000.09 tidak melarang AWS yang bersifat ofensif. Sebaliknya, dinyatakan bahwa jika AWS ofensif memenuhi semua kriteria yang diperlukan, seperti keandalan dalam kondisi realistis, maka pada prinsipnya dapat diotorisasi. Secara eksplisit, Direktif 3000.09 tidak menetapkan batasan apa pun terkait pengembangan dan penggunaan otonomi dalam sistem persenjataan. Departemen Pertahanan AS belum menyetujui senjata otonom apa pun yang berada di luar tiga jenis sistem robotik yang telah diberi 'lampu hijau'. Secara spesifik, Frank Kendall, mantan kepala akuisisi Pentagon, mengatakan: 'Kami belum memiliki apa pun yang bahkan mendekati kemampuan mematikan secara otonom'. Namun, jika pemerintahan militer AS dihadapkan pada situasi seperti itu, Kendall mengatakan bahwa kekhawatirannya adalah memastikan bahwa senjata tersebut mematuhi hukum perang dan bahwa senjata tersebut memungkinkan 'penilaian manusia yang tepat'.

'... Penilaian Dapat Diimplementasikan Melalui Penggunaan Otomatisasi'

Istilah 'penilaian manusia yang tepat' merupakan landasan kebijakan AS tentang AWS, yang diperkenalkan di awal teks dalam konteks berikut: 'Merupakan kebijakan Departemen Pertahanan (DoD) bahwa AWS otonom dan semi-AWS harus dirancang untuk memungkinkan komandan dan operator menerapkan tingkat penilaian manusia yang tepat atas penggunaan kekuatan'. Konsep 'penilaian manusia' tampaknya serupa dengan konsep 'kendali manusia', yang telah dimasukkan dalam laporan yang dibahas secara luas oleh Campaign to Stop Killer Robots (Kampanye), yang diterbitkan hanya dua hari sebelum Direktif 3000.09.. Kampanye tersebut dikenal sebagai pendukung vokal pelarangan AWS. Laporan tersebut menyatakan:

Oleh karena itu, manusia harus tetap memegang kendali atas pilihan untuk menggunakan kekuatan mematikan. Menghilangkan intervensi manusia dalam pilihan penggunaan kekuatan mematikan dapat meningkatkan korban sipil dalam konflik bersenjata.

Dan kemudian dalam ringkasannya:

Saat ini, pejabat militer umumnya mengatakan bahwa manusia akan mempertahankan beberapa tingkat pengawasan atas keputusan penggunaan kekuatan mematikan, tetapi pernyataan mereka sering kali membuka kemungkinan bahwa robot suatu hari nanti dapat memiliki kemampuan untuk membuat pilihan tersebut dengan kekuatan mereka sendiri.

Menariknya, konsep pemerintah AS tentang 'tingkat penilaian manusia yang tepat' atas penggunaan AWS terkadang dianggap sebagai contoh hukum pertama dari konsep 'kendali manusia'. Namun, delegasi AS menjauhkan diri dari deskripsi semacam itu karena penggunaan kata 'kendali'. Mereka berpendapat bahwa penilaian manusia 'atas penggunaan kekuatan' berbeda dari 'kendali manusia' atas senjata. Mereka memberikan contoh berikut:

seorang operator mungkin dapat melakukan kendali yang berarti atas setiap aspek sistem persenjataan, tetapi jika operator hanya secara refleks menekan tombol untuk menyetujui serangan yang direkomendasikan oleh sistem persenjataan, operator tersebut hanya akan melakukan sedikit, jika ada, penilaian atas penggunaan kekuatan. Di sisi lain, penilaian dapat diimplementasikan melalui penggunaan otomatisasi. Misalnya, otomatisasi fungsi yang ekstensif dalam sistem persenjataan dapat memungkinkan operator untuk melakukan penilaian yang lebih baik atas penggunaan kekuatan dengan menghilangkan kebutuhan untuk berfokus pada tugas-tugas dasar dan memberinya lebih banyak waktu untuk memahami situasi yang lebih luas. Demikian pula, penggunaan algoritma atau bahkan fungsi otonom yang mengambil alih kendali dari operator manusia dapat lebih memengaruhi niat manusia dan menghindari kecelakaan.

Delegasi AS menjauhkan diri dari gagasan kendali manusia karena mereka merasa bahwa membingkai perdebatan tentang penggunaan AWS pada 'kendali' terlalu membatasi dan mungkin menyiratkan apa yang disebut kendali manusia langsung. Manusia secara historis menjalankan 'kendali langsung' atas senjata karena senjata telah dilihat hanya sebagai alat di

tangan para pejuang. Dalam arti tertentu, manusia adalah 'penguasa' atas senjata mereka. Konsep 'kendali langsung' senjata oleh manusia dapat ditelusuri ke Konvensi Jenewa 1949 dan Protokol Tambahan 1977 mereka, yang ketentuannya memunculkan gagasan bahwa tanpa kendali atau penggunaan manusia, senjata hanyalah alat belaka. Sebagai contoh, dalam konflik bersenjata, berpartisipasi dalam permusuhan ditunjukkan dengan 'membawa senjata'. Dengan demikian, orang-orang 'yang telah meletakkan senjata mereka' dianggap 'tidak mengambil bagian aktif dalam permusuhan'. Interpretasi LOAC seperti itu akan menyarankan bahwa semua jenis senjata harus dipandu oleh 'kendali langsung' agar sepenuhnya mematuhi hukum. Pandangan ini tercermin oleh ICRC dan HRW, yang berpendapat bahwa tanpa memiliki kendali manusia langsung, AWS akan melanggar prinsip proporsionalitas, antara lain. Delegasi AS tidak setuju dengan argumen ini, dan mereka mengutip berbagai contoh untuk mendukung pemahaman yang lebih luas tentang faktor manusia di medan perang. Salah satu contohnya adalah Sistem Penghindaran Tabrakan Darat Otomatis yang dikembangkan oleh Angkatan Udara AS (USAF) yang telah membantu mencegah apa yang disebut kecelakaan 'penerbangan terkendali ke medan'. Sistem ini mengambil alih kendali pesawat ketika tabrakan yang akan terjadi dengan tanah terdeteksi, dan mengembalikan kendali kepada pilot ketika tabrakan berhasil dihindari. Oleh karena itu, mereka lebih memilih untuk menekankan pada komandan atau operator manusia dan kapasitas mereka untuk menilai kemungkinan efek penggunaan AWS, daripada pada gagasan 'kendali', yang seringkali terbatas dan ditafsirkan terlalu kaku.

Mari kita soroti perbedaan antara kendali manusia dan penilaian manusia dalam konteks yang paling radikal, yaitu, dalam konteks AWS *ofensif*, senjata yang berada di luar tiga jenis sistem robotik yang mendapatkan 'lampu hijau' untuk disetujui dalam Direktif 3000.09. Menurut persyaratan kendali manusia, manusia harus selalu memegang kendali atas keputusan hidup dan mati; dengan demikian, ini berarti bahwa pendelegasian wewenang mematikan kepada mesin untuk membuatnya sendiri tidak diperbolehkan. Hal ini berbeda dengan persyaratan penilaian manusia, di mana pengembangan senjata semacam itu, meskipun tunduk pada pengawasan tambahan, pada prinsipnya diperbolehkan. Perlu ditekankan bagaimana perbedaan semantik yang begitu halus memainkan peran transformatif. Dengan menggunakan kata 'penilaian', Direktif 3000.09 mengalihkan fokus dari kendali langsung pada tingkat keterlibatan dan penargetan ke persyaratan desain senjata yang memungkinkan manusia untuk membuat keputusan yang terinformasi tentang potensi penyebaran senjata. Apresiasi terhadap persyaratan desain menghasilkan dua kewajiban positif berikutnya: (1) bahwa manusia yang menyebarkan sistem harus memahami bagaimana senjata beroperasi di lingkungan yang realistis sehingga manusia dapat membuat keputusan yang terinformasi mengenai penggunaannya, dan (2) untuk memenuhi kewajiban ini, AWS memerlukan tingkat pengujian operasional, verifikasi, validasi, dan evaluasi yang memadai. Namun, persyaratan yang berorientasi pada desain tidak menghasilkan kewajiban untuk tidak mendelegasikan wewenang mematikan kepada mesin untuk membuatnya sendiri selama dua kewajiban positif terpenuhi. Sebaliknya, persyaratan kendali manusia secara eksplisit menyatakan bahwa 'manusia harus mempertahankan kendali atas pilihan untuk menggunakan kekuatan mematikan'.

Untuk menyatakan kembali, kedua kebijakan – konsep kendali manusia dan penilaian manusia – mengakui tantangan yang menimbulkan otonomi mematikan, namun keduanya datang dengan proposisi yang berbeda. Kedua pendekatan tersebut menegaskan bahwa pengembangan senjata robotik telah mencapai titik di mana militer mampu mendelegasikan wewenang mematikan kepada mesin, dan hal ini merupakan tantangan baru bagi para pembuat kebijakan. Namun, satu pendekatan menganjurkan pelarangan senjata yang membuat keputusan penargetannya sendiri, sementara pendekatan kedua membuka peluang bagi potensi pengembangan senjata semacam itu. Apakah ini berarti bahwa baik Kampanye maupun Departemen Pertahanan mengidentifikasi masalah yang sama, tetapi hanya solusinya yang berbeda? Untuk menyelidiki hal ini lebih lanjut, kita harus mempertimbangkan praanggapan apa yang mendasari kedua representasi 'masalah' tersebut.

13.3 ASUMSI KEBIJAKAN PERTAHANAN AS TERKAIT AWS

Istilah 'praanggapan' mengacu pada 'pengetahuan' latar belakang, yang biasanya dianggap sebagai asumsi yang sudah pasti dan bukan asumsi yang dipertanyakan. Penting untuk dicatat bahwa analisis ini tidak bertujuan untuk mengungkap asumsi atau keyakinan yang dipegang oleh para pembuat kebijakan atau mengidentifikasi bias mereka. Sebaliknya, tujuannya adalah untuk mengungkap 'praanggapan yang mengakar kuat dalam representasi masalah'. Untuk melakukannya, terkadang seseorang harus menyelami lebih dalam dari sekadar wacana publik dan juga menyelidiki pendapat para arsitek utama Direktif 3000.09.

Dapat dikatakan bahwa terdapat tiga asumsi inti yang mendasari konstruksi Departemen Pertahanan AS tentang AWS sebagai masalah kebijakan. Pertama, perwakilan Departemen Pertahanan tidak menetapkan batasan tegas apa pun terkait pengembangan otonomi senjata. Mereka membenarkan pendekatan ini dengan mengacu pada prinsip pembalasan yang setara – pemerintah AS tidak ingin berada dalam posisi di mana musuh mereka mengembangkan AWS ofensif yang lebih unggul daripada persenjataan AS. Kedua, otoritas Departemen Pertahanan menetapkan batas potensial pengembangan AWS, jika memang ada, sangat rendah dibandingkan dengan organisasi lain, khususnya Kampanye. Batasan potensial pengembangannya adalah kecerdasan umum buatan (AGI). Ketiga, Departemen Pertahanan tidak menganggap AWS ofensif sebagai senjata yang secara otomatis menimbulkan risiko berlebihan, tetapi justru berpendapat bahwa risiko apa pun dapat dikurangi dengan 'solusi teknis'.

Menyeimbangkan Keselamatan Operasional Dengan Keunggulan Tempur Asimetris

Direktif 3000.09 tidak menjelaskan mengapa Departemen Pertahanan AS membiarkan pintu terbuka untuk membangun dan menyebarkan AWS ofensif. Namun, informasi ini terdapat dalam dokumen pemerintah lainnya. *Strategi Keamanan Nasional* terbaru adalah yang pertama dalam sejarah yang secara khusus menekankan pentingnya AI dan teknologi otonom bagi masa depan militer Amerika. Oleh karena itu, *Strategi Pertahanan Nasional* Departemen Pertahanan berkomitmen untuk 'berinvestasi secara luas dalam penerapan otonomi, AI, dan pembelajaran mesin di bidang militer (...) untuk mendapatkan keunggulan militer yang kompetitif'. Dalam Permintaan Anggaran Tahun Anggaran 2019, AI dan sistem

otonom ditetapkan sebagai prioritas litbang administrasi. AWS secara khusus dianggap sebagai elemen vital dari 'Strategi Offset Ketiga' Departemen Pertahanan. Tujuan dari Strategi Offset Ketiga adalah untuk memastikan keunggulan tempur asimetris yang berkelanjutan bagi AS. Dalam beberapa tahun terakhir, para pesaing AS telah mengembangkan teknologi inovatif yang dapat merusak kepentingan keamanan AS. Pertimbangkan, misalnya, teknologi otonomi kolaboratif seperti swarm. Otonomi kolaboratif adalah jenis kerja sama mesin-mesin di mana beberapa sistem mampu mengoordinasikan tindakan mereka untuk mencapai tujuan bersama. 'Swarm' adalah salah satu kerja sama yang mencakup sekelompok sistem udara homogen dan kecil, yang beroperasi secara kolektif sebagai satu kesatuan yang koheren. Proyek penelitian dan pengembangan di Tiongkok telah menunjukkan bahwa pengembangan algoritma kontrol yang dapat diterapkan untuk berbagai jenis misi swarm dapat dilakukan.

Meskipun Direktif 3000.09 (2012) diumumkan sebelum Strategi Offset Ketiga (2014), hal tersebut tetap memberikan wawasan yang bermanfaat tentang pemikiran awal Departemen Pertahanan (DoD) tentang penggunaan otonomi dalam sistem persenjataan. Kata-kata dalam Direktif tersebut telah dipilih dengan cermat untuk menyeimbangkan pertimbangan keselamatan tambahan demi peningkatan otonomi dalam sistem persenjataan, dengan kurangnya batasan yang kaku terkait pengembangan lebih lanjut otonomi militer untuk mencerminkan tujuan yang kemudian ditetapkan oleh Strategi Offset Ketiga. Contohnya adalah landasan kebijakan yang merupakan *panduan* tingkat penilaian manusia yang tepat atas penggunaan AWS. Kata 'panduan' mungkin membuat pembaca berpikir bahwa Direktif 3000.09 tidak menetapkan kewajiban hukum baru bagi administrasi militer AS. Bahkan, menurut Undang-Undang Prosedur Administratif AS (APA), pernyataan kebijakan dianggap sebagai 'aturan non-legislatif', yang berarti bahwa pernyataan tersebut termasuk dalam definisi 'aturan' tetapi tidak diwajibkan untuk diumumkan melalui prosedur pembuatan peraturan legislatif. Dengan demikian, pernyataan kebijakan tidak memiliki kekuatan hukum yang serupa dengan aturan interpretatif. Yang membedakan aturan interpretatif dari pernyataan kebijakan adalah bahwa yang pertama adalah aturan yang dikeluarkan untuk mengklarifikasi atau menjelaskan undang-undang yang ada, yang terakhir dikeluarkan untuk 'memberi tahu publik tentang cara badan tersebut mengusulkan untuk menjalankan kekuasaan diskresioner'. Aturan legislatif, bertentangan dengan aturan interpretatif dan pernyataan kebijakan, memiliki 'kekuatan dan efek hukum' dan dapat diumumkan hanya jika mereka melalui prosedur pemberitahuan dan komentar publik - sebuah proses di mana publik diberi kesempatan untuk mengomentari versi aturan yang diusulkan dan badan tersebut menanggapi komentar tersebut. Namun, aturan yang melibatkan 'fungsi militer' dikecualikan dari prosedur pemberitahuan dan komentar APA, yang terkadang menyulitkan untuk menentukan apakah aturan tertentu merupakan undang-undang baru, mengklarifikasi undang-undang yang ada atau hanya nasihat tentang pelaksanaan kekuasaan diskresioner suatu badan. Pengadilan di AS berfokus pada bahasa tertentu yang digunakan dalam dokumen saat membuat penentuan ini. Misalnya, bahasa wajib yang menggambarkan kewajiban suatu lembaga dalam pernyataan kebijakan dapat berfungsi sebagai bukti kuat adanya niat untuk mengikat lembaga itu sendiri. Pernyataan lembaga yang 'disusun dalam hal perintah' dapat

dibaca untuk menghilangkan kebijaksanaan lembaga ketika menerapkan kebijakan, mengubah pernyataan tersebut menjadi aturan legislatif. Mengikuti pendekatan ini, jika 'apa yang disebut pernyataan kebijakan dalam tujuan atau kemungkinan efeknya secara sempit membatasi kebijaksanaan administratif, itu akan dianggap apa adanya—aturan hukum substantif yang mengikat'. Menariknya, Direktif 3000.09 memang mengandung beberapa bahasa wajib. Misalnya, mereka mengharuskan AWS ofensif untuk melalui proses peninjauan terperinci sebelum pengembangan dan penerjunan. Namun, tidak pasti bagaimana menafsirkan masalah utama yang dipertaruhkan, yang merupakan konsep penilaian manusia. Direktif 3000.09 menyatakan bahwa AWS 'harus dirancang (...) untuk menjalankan tingkat pertimbangan manusia yang sesuai atas penggunaan kekuatan'. Meskipun berbagai bagian dari Kode Peraturan Federal (CFR) yang mengatur departemen federal menggunakan kata 'harus' untuk menetapkan persyaratan wajib, Mahkamah Agung AS memutuskan bahwa 'harus' juga dapat berarti 'boleh'. Dengan demikian, rumusan Direktif 3000.09 tidak serta-merta menyiratkan apakah persyaratan penilaian manusia atas penggunaan AWS merupakan aturan legislatif atau kebijakan lunak yang memberikan keleluasaan luas kepada departemen militer AS. Mengingat topik AWS masih merupakan bidang penelitian yang baru dan belum pasti bagaimana cara mengatur senjata semacam itu, motivasi di balik ambiguitas ini bisa jadi disengaja dan strategis. Hal ini karena kebijakan penilaian manusia berupaya menyeimbangkan kepentingan yang terkadang saling bertentangan – memastikan keselamatan operasional senjata yang masih relatif sedikit kita ketahui namun tidak membatasi potensi penggunaannya untuk keuntungan tempur. Kebijakan pengendalian manusia, di sisi lain, telah mengabaikan representasi masalah persaingan militer antarnegara dan keinginan untuk mencapai keuntungan tempur asimetris.

Batas Potensi Pengembangan AWS

Direktif 3000.09 tidak menetapkan 'batasan' spesifik apa pun untuk pengembangan AWS. Namun, Robert Work, mantan Wakil Menteri Pertahanan, berpendapat bahwa garis merah pengembangan AI di militer AS *dapat terletak* pada pengembangan kecerdasan umum buatan. 'Bahayanya adalah jika Anda mendapatkan sistem AI umum dan dapat *menulis ulang kodenya sendiri* [MF menekankan] (...) Kami tidak melihat akan pernah menempatkan kekuatan AI sebanyak itu ke dalam senjata apa pun'. Yang mengatakan, Work mengakui bahwa ketika negara lain mulai menggunakan AI umum di medan perang, militer AS mungkin perlu memikirkan kembali pendekatannya. 'Satu-satunya cara kita akan menempuh jalan itu adalah jika ternyata musuh kita melakukannya dan ternyata kita berada pada posisi yang kurang menguntungkan secara operasional (...)'. Pendapat Work tidak terisolasi di antara pejabat tinggi DoD. Dana Deasy, Kepala Informasi Departemen Pertahanan (DoD), mengomentari strategi AI Departemen Pertahanan (DoD AI) yang baru, mengatakan bahwa tema paling sentral di seluruh dokumen adalah 'kebutuhan untuk meningkatkan kecepatan dan kelincahan, yang merupakan cara kami menghadirkan kapabilitas AI di *setiap misi Departemen Pertahanan* [MF menekankan]'.

Kata-kata perwakilan Departemen Pertahanan membantu kita untuk menekankan perbedaan antara asumsi mengenai pengembangan AWS dalam dua kebijakan – satu

mendukung gagasan kendali manusia dan yang lainnya mendukung persyaratan penilaian manusia. Menurut kebijakan kendali manusia, masalah sebenarnya hanyalah fenomena AWS ofensif, bukan otonomi mematikan AGI. Itu adalah perbedaan yang mencolok. AWS ofensif, seperti yang dibahas sebelumnya, adalah fenomena pendelegasian mesin otonom untuk membuat keputusan ofensif yang mematikan. Ini adalah fenomena yang sudah ada; dengan kata lain, sudah ada senjata yang mampu membuat keputusan memamatkannya sendiri. Penelitian ini tidak merujuk di sini pada tiga jenis senjata yang menerima 'lampu hijau' dari Direktif 3000.09, yaitu senjata semi-otonom, senjata otonom yang diawasi secara defensif, dan senjata otonom non-mematikan, non-kinetik – semuanya digunakan saat ini. Ini mengacu pada senjata yang berada di luar tiga kategori dan menggunakan kemampuan otonomi dengan cara yang baru – senjata yang telah digunakan sesekali sebelumnya dan berpotensi digunakan saat ini. Senjata seperti amunisi loitering canggih, yaitu Harpy Israel, di mana tidak ada manusia yang menyetujui target spesifik sebelum keterlibatan. Senjata Harpy ada di gudang senjata berbagai negara saat ini, termasuk Cina, India, Korea Selatan, Chili, dan Turki. Dilaporkan juga bahwa Cina telah merekayasa ulang varian mereka sendiri. Amunisi loitering Harpy digunakan beberapa kali pada tahun 2018 dan 2019 oleh Pasukan Pertahanan Israel untuk menghancurkan baterai SAM Pantsir-S1 Suriah. Departemen Pertahanan AS saat ini memiliki amunisi loitering miniatur berpresisi tinggi yang disebut AeroVironment Switchblade. Switchblade digunakan oleh Angkatan Darat AS di Afghanistan untuk menargetkan 'target bernilai tinggi', seperti pemimpin pemberontak, tim mortir, atau pemberontak yang bepergian dengan kendaraan. Switchblade masih menghubungkan manusia melalui tautan radio yang berfungsi untuk menyetujui target sebelum serangan, menjadikannya senjata semi-otonom, tetapi berpotensi dapat digunakan tanpa intervensi manusia secara langsung. Contoh yang lebih kontroversial dari pendelegasian pengambilan keputusan kepada mesin untuk membuat keputusan yang mematikan adalah Rudal Anti-Kapal Jarak Jauh AGM-158C (LRASM). LRASM adalah rudal jelajah anti-kapal siluman yang dirilis oleh Lockheed Martin dan didasarkan pada Rudal Standoff Udara-ke-Permukaan Gabungan (AGM-158B JASSM-ER) yang menggabungkan sensor frekuensi radio multi-mode, tautan data dan altimeter senjata baru, serta sistem daya yang ditingkatkan. Hal ini memungkinkan LRASM memiliki tingkat kemandirian yang signifikan dari operator manusia. LRASM awalnya diarahkan oleh pilot, tetapi kemudian di tengah perjalanan menuju tujuannya, ia memutuskan komunikasi dengan operatornya. Senjata itu sendiri memutuskan target mana yang dipilih untuk diserang di antara kumpulan yang diberikan. Dengan demikian, ia mampu memilih targetnya sendiri berdasarkan pemrosesan informasi dan aliran data yang berkelanjutan. Senjata LRASM sudah tersedia untuk militer AS, sementara pada Februari 2020, Departemen Luar Negeri AS menyetujui kemungkinan penjualan LRASM ke Australia. Baik Harpy Israel maupun LRASM didasarkan pada kemampuan otonomi canggih; namun, keduanya jauh dari apa yang dianggap sebagai AGI.

AGI merupakan konsep yang berbeda dari AWS yang saat ini digunakan. Meskipun istilah AGI belum memiliki definisi yang diterima secara umum, terdapat beberapa karakteristik yang digunakan secara luas. Pertama, AGI mengacu pada potensi *teknologi masa depan* yang mungkin jauh melampaui kemampuan manusia di semua dimensi. Kedua, AGI

dirancang untuk menampilkan kecerdasan yang tidak terikat pada serangkaian tugas yang sangat spesifik. Dengan kata lain, AGI menggeneralisasi apa yang telah dipelajarinya, termasuk generalisasi ke konteks yang secara kualitatif berbeda dari yang pernah dialaminya sebelumnya atau secara umum menafsirkan tugasnya dalam konteks dunia secara umum. Sistem rekayasa ini belum tersedia secara teknis, dan kedatangannya belum diterima secara universal oleh komunitas AI yang lebih luas.

Terlepas dari kelayakan praktis untuk mencapai AGI, hubungan antara pengembangan AI yang ada dalam sistem persenjataan dan AGI menyoroti representasi masalah yang berbeda dalam dua kubu yang mendukung kendali manusia atau penilaian manusia. Otoritas Departemen Pertahanan menetapkan batas pengembangan potensial AWS sangat rendah dibandingkan dengan Kampanye. Tunduk pada tindakan pencegahan regulasi, pada prinsipnya mereka terbuka untuk mendelegasikan mesin untuk membuat keputusan ofensif dan mematikan. Hal ini secara langsung dikonfirmasi oleh delegasi AS selama pertemuan UN GGE:

Persyaratan DoD Directive 3000.09 bahwa senjata harus dirancang untuk memungkinkan komandan dan operator untuk menggunakan tingkat penilaian manusia yang tepat atas penggunaan kekuatan mencerminkan keputusan yang disengaja untuk mengizinkan senjata yang diprogram untuk membuat 'keputusan' yang berkaitan dengan penargetan.

Selain itu, beberapa arsitek utama DoD Directive 3000.09, seperti Work, melangkah lebih jauh dengan berargumen bahwa jika musuh mulai mengembangkan AGI, AS akan berada pada posisi yang kurang menguntungkan secara operasional dan DoD perlu memikirkan kembali posisinya. Ini berarti bahwa para arsitek Directive 3000.09 membuka kemungkinan untuk mengembangkan dan menggunakan AGI jika ada 'kebutuhan militer yang mendesak'.

Ini bukan berarti bahwa tujuan DoD adalah untuk mengembangkan dan menyebarkan AGI militer. Tak ada strategi resmi yang membangkitkan konsep ini, dan bahkan di antara para arsitek Direktif 3000.09 ada tingkat skeptisisme terhadap kemajuan semacam itu. Namun, tampaknya DoD membiarkan pintu terbuka untuk peningkatan signifikan lebih lanjut dari aplikasi AI yang sudah ada, termasuk dalam fase penargetan dan keterlibatan, dan bahwa persaingan global antarnegara semakin mempercepat proses ini, yang saat ini tidak memiliki garis merah yang keras.

Senjata Yang Menimbulkan Risiko Berlebihan

Asumsi kunci ketiga Departemen Pertahanan di balik kebijakan penilaian manusia mereka adalah bahwa analisis teknis risiko harus menjadi pusat perdebatan tentang potensi pembatasan AWS. Tantangan hukum potensial tunduk pada penilaian teknis murni. Secara khusus, dua isu menonjol ketika membahas faktor risiko AWS: (i) apakah AWS dapat memenuhi prinsip pembedaan LOAC dan (ii) apakah AWS menciptakan 'kesenjangan tanggung jawab'.

Prinsip pembedaan LOAC melarang senjata yang tidak dapat diarahkan pada objek militer tertentu atau jika dampaknya tidak dapat dibatasi sebagaimana disyaratkan oleh hukum internasional, dan jika hasil dalam kedua kasus tersebut adalah bahwa sifat senjata tersebut adalah untuk menyerang target tanpa pembedaan. Menurut Sharkey, AWS yang ada kekurangan tiga komponen teknologi yang diperlukan untuk membedakan antara target yang

sah dan tidak sah. Persyaratan ini adalah: (1) sistem pemrosesan sensorik yang memadai untuk membedakan antara kombatan dan warga sipil; (2) bahasa pemrograman untuk mendefinisikan non-kombatan atau orang yang *hors de combat*; dan (3) kesadaran medan perang atau penalaran akal sehat untuk membantu dalam keputusan diskriminasi. Oleh karena itu, AWS yang ada tidak dapat sesuai dengan prinsip pembedaan dalam situasi konflik yang kompleks di mana mereka akan menghadapi pilihan untuk mengidentifikasi dan terlibat dengan target yang sah yang tersembunyi di antara penduduk sipil.

Pakar militer AS menunjukkan bahwa banyak senjata yang ada, apalagi Harpy, tidak akan memenuhi persyaratan Sharkey dan dengan demikian perlu dihilangkan dari persenjataan modern. Departemen Pertahanan setuju bahwa semua senjata harus mematuhi prinsip-prinsip LOAC, termasuk dengan prinsip pembedaan, tetapi mereka berpendapat bahwa AWS yang ada tidak 'pada dasarnya tidak pandang bulu' dan pada kenyataannya mereka sering meningkatkan kepatuhan terhadap hukum daripada melanggarnya. Bagi Departemen Pertahanan, orang-oranglah yang harus mematuhi LOAC dengan menggunakan senjata dengan cara yang diskriminatif dan proporsional. Misalnya, bahkan jika senjata secara otonom memilih dan menyerang target, 'penggunaannya akan dicegah ketika diperkirakan akan mengakibatkan kerugian insidental pada warga sipil atau objek sipil yang berlebihan dalam kaitannya dengan keuntungan *militer konkret dan langsung yang diharapkan diperoleh*'. Dengan demikian, DoD menggambarkan AWS dari senjata 'yang secara inheren tidak pandang bulu' seperti ranjau atau munisi tandon. Ini berarti bahwa AWS tidak boleh secara otomatis dianggap melanggar hukum. Seseorang tidak boleh mencampuradukkan larangan senjata yang tidak pandang bulu - karena tidak dapat diarahkan ke target yang sah - dengan larangan penggunaan senjata diskriminatif dengan cara yang tidak pandang bulu. Faktanya, AWS yang ada pada dasarnya dirancang untuk bersifat diskriminatif. Mereka hanya dapat menyerang target yang ditunjuk secara khusus yang memenuhi kriteria yang ditentukan oleh algoritma dari antara target yang telah ditentukan sebelumnya.

Sudah ada contoh penggunaan AWS tanpa melanggar prinsip LOAC. Ada situasi di mana AWS dapat memenuhi aturan ini dengan tingkat kemampuan yang cukup rendah untuk membedakan antara target sipil dan militer. Pertama, dalam operasi dalam keadaan yang terdefinisi dengan baik 'tanpa menempatkan warga sipil pada risiko yang berlebihan'. Contoh situasi seperti itu adalah medan perang di mana kombatan dipisahkan secara ketat dari warga sipil dan hanya menempati wilayah tertentu. Kedua, AWS yang ada mungkin hanya menargetkan senjata musuh, bukan individu yang mengoperasikannya, hingga individu tersebut menimbulkan potensi ancaman. Terakhir, senjata semacam itu mungkin juga beroperasi di tempat yang tidak ada warga sipil. Dalam skenario seperti itu, AWS hanya akan memilih dan menyerang mesin musuh lainnya. Berdasarkan pertimbangan di atas, setidaknya ada tiga rangkaian keadaan potensial di mana AWS yang ada mungkin dapat mematuhi prinsip pembedaan.

Ada dua poin penting dalam argumentasi DoD. Pertama, AWS yang ada tidak 'pada dasarnya tidak pandang bulu'. Kedua, meskipun AWS yang sedang berkembang saat ini mungkin menimbulkan tantangan tertentu terkait kemampuan teknis untuk mendiskriminasi

warga sipil, peningkatan lebih lanjut dalam teknologi akan 'memperbaiki' tantangan ini. Gagasan 'memperbaiki teknologi' ini digaungkan oleh pernyataan perwakilan AS selama pertemuan GGE di Jenewa:

Teknologi yang sedang berkembang sulit diatur karena teknologi terus berubah seiring perkembangan para ilmuwan dan insinyur. Praktik terbaik saat ini mungkin bukan praktik terbaik dalam waktu dekat. Demikian pula, sistem persenjataan yang, jika dibangun saat ini, berisiko menciptakan efek yang tidak pandang bulu, mungkin, jika dibangun dengan teknologi masa depan, terbukti lebih diskriminatif daripada alternatif yang ada dengan mengurangi risiko korban sipil.

Menariknya, pandangan Kampanye tentang 'perbaikan teknologi' di masa depan atas kekurangan senjata yang ada kurang jelas. Dalam laporan *Losing Humanity*, para penulis Kampanye berpendapat di satu tempat bahwa AGI, 'yang akan mencoba meniru pemikiran manusia', merupakan opsi potensial untuk mendorong kepatuhan senjata otonom terhadap LOAC. Lebih lanjut, dalam laporan tersebut, mereka tampaknya meragukan apakah teknologi benar-benar dapat meniru tindakan manusia:

Sekalipun pengembangan senjata yang sepenuhnya otonom dengan kognisi seperti manusia menjadi layak, senjata tersebut akan kekurangan kualitas manusia tertentu, seperti emosi, kasih sayang, dan kemampuan untuk memahami manusia. Akibatnya, adopsi senjata semacam itu secara luas masih akan menimbulkan masalah hukum yang meresahkan dan menimbulkan ancaman lain bagi warga sipil.

Namun, untuk bergerak melampaui perdebatan tentang kelayakan teknis merancang senjata agar sesuai dengan LOAC, Kampanye beralih ke masalah penting lainnya – masalah tanggung jawab atas kesalahan senjata.

Masalah Kesenjangan Tanggung Jawab Kembali Mengalihkan Perhatian Ke Analisis Risiko

Potensial, mungkin ada situasi di mana tidak ada yang bertanggung jawab atas kesalahan mesin. Inilah yang disebut 'kesenjangan tanggung jawab'. Beberapa berpendapat bahwa semakin otonom sistem senjata, semakin kecil kemungkinan untuk meminta pertanggungjawaban dari perancang atau penggunaannya atas tindakan mereka. Namun, ketidakmungkinan menghukum agen buatan berarti seseorang tidak dapat meminta pertanggungjawaban mesin. Kesenjangan tanggung jawab muncul ketika dalam pelaksanaan keputusan penargetan, AWS pasti telah melakukan sesuatu yang tidak diprogramkan langsung oleh operator, sehingga operator tidak dapat disalahkan atas penyalahgunaan kekuatan. Di satu sisi, elemen ketidakpastian ini sering kali menjamin fleksibilitas penyesuaian mesin terhadap lingkungan yang dinamis. Para pemrogram sengaja merancang sistem ini agar dapat merespons perubahan secara langsung, alih-alih mengantisipasi setiap kemungkinan yang mungkin muncul. Di sisi lain, unsur ketidakpastian ini tampaknya menjadi masalah khusus dalam penerapan kekuatan terhadap target. Masalahnya adalah AWS sengaja menjauhkan operator dari penegakan keputusan penargetan. Operator muncul pada tahap awal rantai

kausal yang mengarah pada penerapan kekuatan terhadap target, tetapi bukan pada tahap akhir.

Masalah kesenjangan tanggung jawab dapat dibingkai sebagai masalah hukum. Dapat dikatakan bahwa mekanisme tanggung jawab hukum yang ada tidak sesuai dan tidak memadai untuk sepenuhnya mengatasi kerugian yang melanggar hukum yang mungkin disebabkan oleh AWS dan akibatnya manusia yang terlibat dalam produksi atau penggunaan AWS akan 'lolos dari tanggung jawab atas penderitaan yang disebabkan oleh senjata tersebut'. Perwakilan DoD terlibat dengan argumen hukum, tetapi mereka mempersempitnya menjadi pertimbangan teknis murni mengenai analisis risiko. Proses ini terjadi dalam langkah-langkah berikut.

Strategi awal DoD adalah mengecilkan relevansi 'kesenjangan tanggung jawab'. Mereka berpendapat bahwa 'kesenjangan tanggung jawab' tidak terjadi karena, pada prinsipnya, LOAC terutama berurusan dengan negara, bukan individu. Tidak seperti tanggung jawab pidana individu, tanggung jawab negara tidak didasarkan pada konsep kesalahan pribadi, tetapi pada atribusi kepada negara atas (kesalahan) tindakan organ atau agennya. Tanggung jawab pidana individu hanya dapat dikaitkan di bawah LOAC pada pelanggaran hukum yang paling serius, seperti kejahatan perang, kejahatan terhadap kemanusiaan, agresi, dan genosida. Namun, Kampanye berpendapat bahwa penggunaan AWS berpotensi melakukan kejahatan tersebut, dan dengan demikian muncul kesenjangan tanggung jawab. Menurut konsep tanggung jawab pidana individu, pelanggaran pidana disebabkan oleh kesengajaan atau kelalaian. Ketika agen buatan secara sengaja diarahkan untuk menyakiti atau menyakiti disebabkan oleh kelalaian maka operator manusia bertanggung jawab secara pidana. Ketika ketiadaan niat (*mens rea*) operator tidak dapat dibuktikan, maka muncullah kesenjangan tanggung jawab.

DoD tidak setuju dan mengklaim bahwa dalam seluruh rantai komando selalu ada unsur manusia yang dapat ditemukan baik di tingkat pemrograman maupun di tingkat operasional. Dengan demikian, pemrogram dapat dianggap bertanggung jawab jika ia memprogram AWS sedemikian rupa sehingga secara sengaja melanggar hukum konflik bersenjata, atau operator dapat bertanggung jawab jika ia memutuskan untuk mengoperasikan AWS dengan cara yang melanggar hukum. Dalam situasi di mana baik pemrogram maupun komandan tidak dapat diidentifikasi pada awalnya, ketidakmungkinan untuk menetapkan tanggung jawab pidana kepada seseorang bukan disebabkan oleh fakta bahwa tindakan berbahaya tersebut dilakukan oleh AWS, tetapi masalahnya adalah ketidakmungkinan untuk *mengumpulkan bukti* yang memungkinkan identifikasi yang tepat dari unsur manusia yang relevan yang bertanggung jawab atas kesalahan mesin.

Dengan demikian, diskusi tentang tanggung jawab dimulai sebagai wacana hukum, tetapi kemudian berkembang menjadi perdebatan tentang ambang batas risiko yang wajar dan standar risiko. Analisis risiko pada gilirannya sebagian besar didominasi oleh pengetahuan teknis tentang kinerja AWS saat ini dan potensi kinerja di masa mendatang, spesifikasi pengujian dan pelatihannya. Di sini, sekali lagi, seperti halnya masalah hukum perbedaan, Departemen Pertahanan dan para ahlinya mengalihkan fokus dari argumen hukum dan moral ke pertimbangan yang hampir sepenuhnya teknis. "Standar kehati-hatian atau perhatian yang diperlukan dalam menjalankan operasi militer terkait perlindungan warga sipil merupakan

pertanyaan kompleks yang tidak dapat dijawab secara sederhana oleh hukum perang", kita membaca. Departemen Pertahanan lebih lanjut menyatakan bahwa standar ini harus dinilai berdasarkan praktik umum negara bagian dan standar umum profesi militer dalam menjalankan operasi. Secara khusus, langkah terbaik untuk mendorong "akuntabilitas" adalah:

pelatihan sistem persenjataan dan pengujian sistem persenjataan yang ketat dapat membantu para komandan mengetahui kemungkinan dampak dari penggunaan sistem persenjataan tersebut.

Dengan mempertimbangkan hal-hal di atas, meskipun Departemen Pertahanan menyadari peningkatan risiko yang terkait dengan penggunaan otonomi saat ini, Departemen Pertahanan tidak serta merta menganggap AWS sebagai senjata yang secara inheren menimbulkan risiko yang terlalu berlebihan. Keterbatasan potensial dari pengembangan militer adalah AGI, meskipun tidak tanpa reservasi. Meskipun demikian, lompatan signifikan dari kemajuan teknis diperlukan untuk menutup kesenjangan antara senjata yang ada yang memanfaatkan AI sempit dan AGI yang dipersenjatai. Ini berarti bahwa militer AS, tunduk pada semua tindakan pencegahan, membiarkan diri mereka terbuka untuk kemajuan lebih lanjut dari AWS, didorong oleh persaingan internasional, yang hanya mempercepat lajunya. DoD mengambil pandangan yang secara kualitatif berbeda tentang masalah hukum diskriminasi dan tanggung jawab yang relatif terhadap Kampanye. Mereka mengalihkan konsekuensi hukum dari masalah tersebut ke masalah teknis dan membangun kepercayaan dalam pandangan mereka berdasarkan kecakapan teknologi mereka dan kemampuan pengujian dan pemantauan yang relevan. Dengan demikian, tantangan hukum dan etika baru yang potensial untuk kontrol manusia-mesin yang timbul dari AWS sebagian besar, jika tidak secara eksklusif, ditundukkan pada keahlian teknis dan 'pengetahuan' militer.

13.4 REPRESENTASI PERTAHANAN AS YANG LUPUT

Dengan beralih dari 'kendali manusia' ke 'penilaian manusia', dalam bagian ini, dapat dikatakan bahwa kebijakan tersebut menantang gagasan tentang peran aktif manusia dalam pengoperasian AWS. Respons terhadap masalah meningkatnya risiko yang terkait dengan penggunaan AWS adalah 'lebih banyak otonomi, lebih sedikit manusiawi', tetapi asumsi ini sebagian besar telah ditinggalkan dari wacana representasi masalah, menggantikan gagasan samar tentang 'penilaian manusia'. Lebih lanjut, Direktif 3000.09 mengecualikan domain siber dari cakupannya dan dengan demikian menciptakan kekosongan hukum dan kebijakan terkait potensi pengembangan dan penggunaan senjata siber otonom ofensif – sebuah ancaman yang bisa dibilang lebih persisten daripada penggunaan insidental AWS kinetik. Akhirnya, masalah yang dirumuskan oleh Departemen Pertahanan terutama dibingkai sebagai masalah teknis. Meskipun kebijakan AS mengakui tantangan spesifik yang terkait dengan pengembangan dan penggunaan AWS, kebijakan tersebut menyisakan pertimbangan yang tidak bermasalah bahwa senjata semacam itu, meskipun memenuhi semua tinjauan teknis yang diperlukan, masih dapat dianggap tidak dapat diterima semata-mata atas dasar etika. Lebih lanjut, ketiga

argumen ini dapat menginformasikan permasalahan alternatif dan representasi kebijakan AWS.

Kendali Desain dan Penghapusan Kendali Manusia

Bagian ini mengidentifikasi tiga transformasi penting, meskipun tidak diakui secara luas, yang muncul antara tahun 2005 dan 2011, dan berpuncak pada Direktif Departemen Pertahanan 3000.09. Semuanya mengarah pada penggantian bertahap gagasan pengendalian manusia dengan konsep penilaian manusia.

Pertama, meskipun strategi Departemen Pertahanan sejak tahun 2005 menekankan peran faktor manusia, seiring waktu, kita dapat mengamati dalam dokumen militer penurunan pentingnya pengendalian manusia secara langsung dan meningkatnya relevansi pengendalian tidak langsung, terutama pada tahap desain. Rencana Penerbangan Angkatan Udara 2009 mengidentifikasi pergeseran peran manusia dari operator aktif menjadi agen pemantau. Rencana tersebut menjelaskan:

Manusia akan semakin tidak lagi berada 'dalam lingkaran', melainkan 'di dalam lingkaran' – memantau pelaksanaan keputusan tertentu. Bersamaan dengan itu, kemajuan dalam AI akan memungkinkan sistem untuk membuat keputusan tempur dan bertindak dalam batasan hukum dan kebijakan tanpa perlu memerlukan masukan manusia.

Kita dapat melihat bagaimana deskripsi sistem berbasis AI mempersiapkan situasi untuk penghapusan keterlibatan manusia secara langsung. Rencana tersebut juga menekankan perubahan peran operator manusia, yang sebelumnya berada di dalam lingkaran menjadi berada di dalam lingkaran. Peran mereka adalah untuk 'memantau pelaksanaan operasi' dan mempertahankan 'kemampuan untuk mengesampingkan sistem selama misi'. Dalam Rencana yang sama, kita membaca tentang penanaman kualitas manusia ke dalam mesin pada tingkat desain: 'pemrograman sistem akan didasarkan pada niat manusia'.

Transformasi penting kedua adalah penghapusan referensi untuk kendali manusia langsung setelah Direktif DoD 3000.09. Pada tahun 2009, AU AS meminta kebijakan untuk memandu pengembangan kemampuan persenjataan masa depan, termasuk sistem yang sepenuhnya otonom. Kantor Kebijakan Pertahanan mulai menerima pertanyaan tentang masalah hukum dan etika terkait penggunaan otonomi mematikan, sementara berbagai cabang militer merespons tanpa koordinasi. Angkatan Darat AS awalnya mengklaim bahwa mereka 'tidak akan pernah mendelegasikan keputusan penggunaan kekuatan kepada robot'. AU AS memiliki perspektif yang berbeda. Setelah dua tahun, pada tahun 2011, Departemen Pertahanan AS merespons dengan Peta Jalan Sistem Tak Berawak sementara yang menyatakan:

Pedoman kebijakan khususnya akan diperlukan untuk sistem otonom yang melibatkan penerapan kekuatan (...) Untuk masa mendatang yang dapat diperkirakan, keputusan tentang penggunaan kekuatan dan pilihan target individu mana yang akan dikenai kekuatan mematikan akan tetap berada di "bawah kendali manusia" [MF menekankan] dalam sistem tak berawak.

Namun, dalam dokumen yang sama, kita membaca bahwa Departemen Pertahanan membayangkan sistem tanpa awak untuk beroperasi dengan sistem berawak 'sambil secara bertahap *mengurangi tingkat kendali manusia* [MF menekankan] dan pengambilan keputusan yang diperlukan untuk bagian tanpa awak dari struktur pasukan'. Dalam Direktif Departemen Pertahanan 3000.09 dari tahun 2012, gagasan kendali manusia digantikan oleh 'penilaian manusia', menandakan pergeseran yang lebih luas menuju kendali pada tingkat desain relatif terhadap kendali manusia langsung. Dalam Peta Jalan Sistem Tanpa Awak dari tahun 2013, gagasan kendali manusia digunakan dengan cara berikut:

Penelitian dan pengembangan dalam otomasi sedang berkembang dari keadaan sistem otomatis yang membutuhkan kendali manusia menuju keadaan sistem otonom yang mampu membuat keputusan dan bereaksi tanpa interaksi manusia.

Para praktisi militer saat ini berpendapat bahwa Direktif DoD 3000.09 bersifat transformatif karena merupakan kebijakan pertama terkait AWS. Namun, revolusi yang paling signifikan tidak disadari – DoD menghapus konsep 'kendali manusia' dan dengan memasukkan frasa 'tingkat penilaian manusia yang sesuai atas penggunaan kekuatan', mereka membuka peluang bagi penerapan kendali manusia secara langsung atas penggunaan AWS, termasuk dalam situasi ofensif.

Aspek ketiga dari perubahan kontrol manusia-mesin antara tahun 2005 dan 2012 terkait dengan pengenalan 'tingkat otonomi' oleh DoD sebagai kerangka kerja untuk menginformasikan interaksi dengan operator manusia. Dalam Peta Jalan 2005, DoD mendefinisikan sepuluh tingkat otonomi mulai dari sistem yang dipandu dari jarak jauh hingga sepenuhnya otonom. Peta Jalan 2009 secara khusus meminta operator dan komandan untuk 'mempertahankan kemampuan untuk *memperbaiki* [MF menekankan] tingkat otonomi'. Dan kemudian: 'Tingkat otonomi harus *disesuaikan* secara dinamis [MF menekankan] berdasarkan beban kerja dan maksud yang dirasakan oleh operator'. Panduan ini mencoba untuk membantu proses pengembangan dan penggunaan senjata dengan mengelompokkan fungsi untuk skenario umum. Mereka mengasumsikan pemisahan tugas antara manusia dan mesin. Manusia diatur untuk merespons situasi yang dinamis dan beralih ke 'mode otomatisasi' yang tepat daripada berinteraksi secara koaktif dengan mesin untuk mencapai kemampuan terbaik. Konsep tingkat otonomi telah dikritik oleh praktisi yang berpendapat bahwa kerangka kerja 'tingkat otonomi' hanya mewakili situasi di mana peningkatan otomatisasi harus disertai dengan lebih sedikit kendali manusia, yang menciptakan trade-off yang tidak perlu dan tidak secara akurat mewakili kompleksitas campur tangan manusia-mesin yang tertanam dalam AWS. Kritik tersebut juga telah dilayangkan di dalam DoD. Beberapa bulan setelah diperkenalkannya Direktif 3000.09, yang menegaskan kembali gagasan tingkat otonomi, DSB menerbitkan laporan yang mengkritik keras konsep ini. DSB berpendapat untuk mengganti 'tingkat otonomi' dengan kerangka kerja yang berfokus pada alokasi eksplisit fungsi dan tanggung jawab kognitif antara manusia dan mesin untuk mencapai kemampuan tertentu.

Namun, terlepas dari kritik-kritik ini, konsep 'tingkat otonomi' masih digunakan dalam militer AS, meskipun laporan DSB menyarankan pendekatan yang berbeda.

Singkatnya, dalam debat akademis, banyak penulis menyamakan konsep kendali manusia dengan gagasan penilaian manusia, meskipun Departemen Pertahanan (DoD) menjauhkan diri dari klaim tersebut. Departemen Pertahanan (DoD) menetapkan Direktif 3000.09 sebagai respons terhadap masalah meningkatnya risiko yang terkait dengan penggunaan AWS, tetapi mendasarkan respons tersebut pada pendekatan yang lebih luas, yaitu 'menghilangkan awak' dan menolak kendali manusia. Direktif 3000.09 bahkan membuka peluang bagi penerapan tanpa penilaian manusia sama sekali. Keyakinan Departemen Pertahanan saat ini tentang 'lebih banyak otonomi, kurang manusiawi' sebagai respons terhadap risiko operasional khususnya bermasalah, tetapi sebagian besar telah diabaikan dalam representasi masalah tersebut. Cara alternatif untuk melihat masalah ini adalah dengan menghargai perlunya faktor manusia dengan menyatakan bahwa semakin canggih suatu senjata, semakin kompleks sistem kontrolnya, dan dengan demikian secara paradoks kontribusi operator manusia lebih penting. Pendekatan alternatif ini dapat berasal dari kerangka kerja AI yang Berpusat pada Manusia, yang mendalilkan untuk merancang teknologi yang menawarkan kontrol manusia tingkat tinggi dan otomatisasi komputer tingkat tinggi. Menurut perspektif ini, terlepas dari tingkat otonomi sistem, manusia harus selalu memiliki kemungkinan untuk mengesampingkan tindakan mesin. Selain itu, manusia juga harus melakukan pengawasan aktif terhadap sistem senjata untuk menghindari 'kesombongan algoritmik', yang setara dengan harapan yang tidak masuk akal mengenai kinerja mesin. Ini hanyalah contoh yang paling umum, tetapi argumen utamanya adalah bahwa ada alternatif untuk formulasi masalah DoD dengan mendelegasikan otonomi mematikan kepada mesin yang mempersulit 'pengembangan teknologi alami' yang tampaknya menuju lebih banyak otonomi dan lebih sedikit keterlibatan manusia.

Pengecualian senjata siber dan kompleksitasnya

Direktif 3000.09 secara eksplisit mengecualikan sistem otonom atau semi-otonom untuk operasi dunia maya. Namun, senjata siber sebenarnya dapat mewakili ancaman yang lebih besar daripada penggunaan AWS kinetik yang terisolasi karena dapat menghasilkan efek berbahaya di seluruh internet. Tantangan paling signifikan datang dari senjata siber ofensif otonom, yang sudah digunakan sejak operasi Stuxnet. Stuxnet bukan hanya senjata siber pertama yang menyebabkan kerusakan fisik, tetapi juga merupakan senjata yang secara otonom melakukan serangannya, meskipun otonominya sebagian besar masih dibatasi. Stuxnet beroperasi dengan cara yang mirip dengan amunisi homing, mencari target tertentu (dalam hal ini untuk mengoperasikan pengontrol logika terprogram yang digunakan dalam aplikasi industri) dan ketika target ditemukan, Stuxnet mengubah pengaturannya dan mengambil kendali. Namun, Stuxnet memiliki sejumlah perlindungan untuk membatasi penyebaran dan efeknya; misalnya, ia tidak dapat memodifikasi dirinya sendiri secara otonom dan memiliki tanggal penghentian sendiri. Meskipun pembaruan perangkat lunak baru umumnya berasal dari manusia, para pakar keamanan telah memperingatkan tentang senjata siber ofensif otonom yang lebih canggih dengan kemampuan untuk mereplikasi diri.

Seperti halnya senjata kinetik, di sini lagi-lagi perwakilan Departemen Pertahanan khawatir bahwa musuh dapat mendorong AS untuk meluncurkan AWS ofensif, meskipun tidak adanya aturan yang jelas mengenai penggunaannya. Kita membaca dalam Strategi Keamanan Nasional bahwa pesaing AS tidak hanya mencapai kemajuan signifikan dalam mengintegrasikan kemampuan analitik data dalam operasi militer, tetapi mereka juga mengembangkan senjata canggih yang dapat mengancam *arsitektur komando dan kendali* AS saat ini. Instruksi komando dan kendali canggih digunakan untuk mengidentifikasi penyerang, menghubungkan serangan, dan menghentikan aktivitas jahat yang sedang berlangsung.

Namun, setidaknya ada empat tantangan operasional untuk memiliki arsitektur komando dan kendali yang efektif dalam domain siber. Pertama, ada masalah ketidakpastian ofensif. Agen malware cerdas dengan kapasitas belajar mandiri dapat mempelajari dan mengesampingkan tindakan defensif untuk mengeksploitasi potensi kerentanan sistem. Tantangan kedua berkaitan dengan ketidakdeteksian pelanggaran, yang disebut dalam studi cyber 'masalah atribusi'. Malware yang kompleks sulit dideteksi dan, bahkan ketika dikenali, seseorang hanya dapat memberikan penjelasan untuk intrusi yang diketahui. Tantangannya adalah untuk menghadapi prospek intrusi permanen infrastruktur pembela ketika skala dan ruang lingkup infiltrasi dapat menimbulkan masalah keamanan yang signifikan. Tantangan ketiga berkaitan dengan kompleksitas sistem pertahanan. Sementara penyerang biasanya hanya perlu memahami prosedur masuk, pembela harus melindungi seluruh jaringan terhadap banyak titik kerentanan yang saling terkait. Terakhir, arsitektur perintah dan kontrol tradisional berada di bawah tekanan dari risiko rantai pasokan, seperti produsen yang memperkenalkan kerentanan ke dalam komponen spesifik suatu sistem.

Yang mengatakan, Direktif 3000.09 mengecualikan domain cyber dan dengan demikian menciptakan kekosongan hukum dan kebijakan ketika datang ke potensi pengembangan dan penggunaan senjata cyber otonom ofensif. Oleh karena itu, cara alternatif untuk melihat masalah yang dirumuskan adalah dengan mempertimbangkan kerangka kerja yang saat ini terpisah untuk domain cyber dan AWS kinetik bersama-sama dan untuk mengeksplorasi apakah kontrol yang ada dari Direktif 3000.09 sesuai dengan operasi senjata cyber otonom. Jika tidak, ada baiknya mengeksplorasi apa yang membuat senjata kinetik istimewa sehingga mereka memerlukan peningkatan pengawasan peraturan. Jawaban potensial adalah bahwa AWS kinetik bisa mematikan, sedangkan senjata cyber tidak. Namun, senjata cyber dapat menyebabkan kerusakan yang mematikan atau terkait sebagai efek samping atau dapat menyebabkan kerusakan yang signifikan, bahkan jika tidak mematikan. Dengan demikian, seseorang dapat meningkatkan wawasan ini, mengungkap kompleksitas terkait senjata siber, dan menantang representasi masalah dalam bentuk saat ini oleh Departemen Pertahanan.

Melampaui Pertimbangan Hukum Dan Teknis – Martabat Manusia Dan Keterlibatan Moral

Hal yang tidak menjadi masalah dalam pendekatan Departemen Pertahanan AS terhadap AWS adalah bahwa senjata semacam itu, meskipun memenuhi semua tinjauan hukum dan teknis yang diperlukan, masih dapat dianggap tidak dapat diterima semata-mata atas dasar etika. Dengan kata lain, dapat dikatakan bahwa mendelegasikan kemampuan

mematikan kepada mesin untuk membuat keputusan pada dasarnya tidak etis. Patut diakui bahwa argumen-argumen semacam ini mulai mendapatkan lebih banyak perhatian sekarang karena perdebatan tentang AWS telah matang. Misalnya, dalam laporan pertama yang diterbitkan oleh Kampanye, kata 'martabat' tidak muncul, sementara pertimbangan etika dibahas terutama dalam konteks potensi kepatuhan terhadap hukum humaniter internasional. Ini bukan berarti bahwa argumen bahwa AWS melanggar martabat manusia, terlepas dari pertimbangan hukum dan teknis, tidak dikemukakan oleh Kampanye. Laporan Kampanye merujuk, misalnya, pada 'dampak mendalam' pada 'impersonalisasi pertempuran' yang dapat menghilangkan beberapa keberatan naluriah terhadap pembunuhan. Namun, terlepas dari ungkapan-ungkapan ini, masih ada seruan dari komunitas akademis untuk lebih mengurangi penekanan pada kerangka hukum dan teknis demi argumen etis yang lebih mendalam bahwa penggunaan AWS akan melanggar martabat manusia.

Persoalan martabat manusia merupakan masalah yang kompleks, tetapi sering dipahami sebagai status khusus yang dikaitkan dengan manusia yang darinya muncul hak dan kewajiban tertentu. Dalam bidang peperangan, pertanyaannya adalah apakah sesuatu yang bernilai moral yang merupakan status khusus manusia hilang ketika manusia digantikan oleh mesin dalam penggunaan kekuatan. Leveringhaus berpendapat bahwa penggantian agensi manusia dengan agensi buatan pada titik penyampaian kekuatan tidak diinginkan secara moral karena mengarah pada pelepasan moral. Argumen ini tidak menunjukkan bahwa operator manusia tidak sepenuhnya terlepas. Mereka harus memastikan bahwa penggunaan AWS tidak menyebabkan risiko yang berlebihan. Namun, mereka tidak *sepenuhnya* terlibat secara moral. Hal ini karena menjadi manusia yang sepenuhnya terlibat secara moral lebih dari sekadar menghormati hak orang lain. Keterlibatan moral berarti bertindak atas dasar alasan yang tidak sepenuhnya berbasis hak, seperti pengakuan atas kemanusiaan bersama, kepedulian terhadap yang rentan, atau rasa kasihan dan belas kasihan. Dengan demikian, penggantian agensi manusia dengan agensi buatan setidaknya mengarah pada pelepasan moral sebagian.

13.5 KESIMPULAN

Departemen Pertahanan AS (DoD) memperkenalkan kebijakan AWS sebagai respons terhadap masalah potensi penggunaan AWS untuk tujuan ofensif, yang meningkatkan risiko operasional akibat yang tidak diinginkan. Direktif 3000.09 memperkenalkan persyaratan penilaian manusia atas penggunaan senjata tersebut untuk memitigasi risiko yang terkait dengan penggunaan senjata tersebut, tetapi dengan menggunakan kata 'penilaian', kebijakan tersebut mengalihkan fokus dari kendali langsung pada tingkat keterlibatan dan penargetan ke persyaratan desain senjata.

Pendekatan Departemen Pertahanan AS didasarkan pada asumsi bahwa pengembangan otonomi mematikan dibenarkan oleh persaingan global antarnegara, dan dengan demikian tidak ada hambatan khusus untuk pengembangan. Dengan demikian, hal ini sesuai dengan kontur yang lebih luas dari Strategi Offset Ketiga yang bertujuan untuk membangun kapabilitas teknologi yang unggul dibandingkan dengan negara-negara seperti Tiongkok dan Rusia. Bertentangan dengan Kampanye, Departemen Pertahanan tidak

menganggap AWS ofensif sebagai senjata yang tentu saja menimbulkan risiko yang terlalu berlebihan. Argumen hukum yang menentang AWS – ketidakmampuan untuk mendiskriminasi warga sipil dan potensi kesenjangan tanggung jawab akibat kesalahan AWS – terutama tunduk pada wacana 'teknis' dan pengetahuan militer.

Dinyatakan bahwa kebijakan Departemen Pertahanan AS tentang AWS menantang gagasan tentang peran aktif manusia dalam pengoperasian senjata semacam itu. Respons mereka terhadap masalah risiko operasional yang terkait dengan penggunaan AWS adalah 'lebih banyak otonomi, lebih sedikit manusiawi', tetapi asumsi ini sebagian besar telah ditinggalkan dari wacana representasi masalah. Dengan demikian, seseorang dapat membangun representasi masalah alternatif yang menghargai perlunya campur tangan manusia secara langsung dengan berargumen bahwa semakin canggih suatu senjata, semakin kompleks sistem kendalinya, sehingga kontribusi operator manusia semakin krusial. Lebih lanjut, Direktif 3000.09 mengecualikan domain siber dari cakupannya dan dengan demikian menciptakan kekosongan hukum dan kebijakan dalam hal pengembangan dan penggunaan senjata siber otonom ofensif. Salah satu representasi permasalahan alternatif adalah dengan mempertimbangkan kerangka kerja yang saat ini terpisah untuk domain siber dan AWS kinetik secara bersamaan, dan untuk mengeksplorasi apakah kontrol yang ada dari Direktif 3000.09 sesuai dengan kontrol pada operasi senjata siber otonom, dan jika tidak, mengapa demikian. Terakhir, permasalahan yang dirumuskan oleh Departemen Pertahanan (DoD) terutama dibingkai sebagai permasalahan teknis. Departemen Pertahanan tidak mempersoalkan pertimbangan bahwa senjata semacam itu, meskipun telah memenuhi semua tinjauan teknis yang diperlukan, masih dapat dianggap tidak dapat diterima secara etis. Representasi permasalahan alternatif dapat menantang pengembangan dan penggunaan AWS karena kekhawatiran tentang martabat manusia dalam peperangan.

13.6 RINGKASAN

Dalam beberapa tahun terakhir, perhatian publik semakin terfokus pada implementasi sistem terkait AI, tanpa analisis hukum dan peraturan yang memadai. Oleh karena itu, publikasi ini bertujuan untuk mengisi kesenjangan ini dengan memandu pembaca tidak hanya melalui berbagai aplikasi AI, tetapi juga melalui aspek hukum yang terkait. Penelitian dalam volume ini mengungkapkan bahwa AI masih memiliki jalan panjang sebelum dapat dianggap aman dan tepercaya. Tingkat pengembangan perangkat AI sangat beragam, dan kekuatan serta kelemahannya telah diungkapkan dalam berbagai bab. Meskipun teknologi memungkinkan penggunaannya untuk melakukan tugas-tugas tertentu lebih cepat dan seringkali dengan cara yang lebih akurat, "skeptisisme" umum terhadap AI di antara berbagai regulator secara konsisten berkisar pada kerentanan utama tertentu, termasuk akuntabilitas, transparansi, kemudahan dijelaskan, dan perlindungan hak-hak asasi. Meskipun demikian, pandemi COVID-19 telah menjadi katalis yang kuat bagi banyak industri untuk merangkul teknologi informasi, komunikasi, dan AI lebih erat. Dengan banyaknya bisnis yang menerapkan atau mengeksplorasi opsi AI, sangat penting bagi perancang, regulator, dan pelaku bisnis AI untuk bekerja sama

membangun dan terus meningkatkan sistem dengan menerapkan kerangka kerja regulasi dan etika yang baik.

Meskipun Uni Eropa memiliki beberapa regulasi privasi dan keamanan terketat di dunia, kerangka hukum umum untuk AI masih dalam tahap pengembangan. Saat ini, belum ada regulasi komprehensif tentang AI di yurisdiksi mana pun. Undang-undang nasional beberapa negara memperkenalkan regulasi untuk penggunaan teknologi tertentu, tetapi, secara umum, aspek hukum pengaturan penggunaan AI masih dalam tahap pembahasan atau pengembangan. Harmonisasi regulasi AI di berbagai yurisdiksi juga masih minim. Kerangka kerja regulasi tertentu yang berlaku untuk industri tertentu mencoba mengimbangi kemajuan teknologi yang pesat dengan menciptakan berbagai hak sui generis. Contoh yang baik dari hal ini adalah sistem hak cipta yang berlaku untuk banyak industri kreatif, seperti musik. Namun, sebagaimana disajikan dalam Bab 7, kerangka kerja tersebut tampaknya tidak menangani beberapa permasalahan mendasar dalam industri tersebut. Segala upaya harmonisasi regulasi AI di masa mendatang harus bertujuan untuk memberikan kepastian hukum yang diperlukan guna memfasilitasi inovasi dan investasi di bidang AI, sekaligus melindungi hak-hak dasar dan memastikan aplikasi AI digunakan dengan aman. Lebih lanjut, upaya tersebut harus menyediakan mekanisme pengawasan dan penegakan hukum yang memungkinkan harmonisasi nasional atau regional, serta penerapan tindakan korektif dan sanksi. Agar kerangka regulasi efektif, dampak dan efektivitasnya juga perlu direvisi secara sistematis. Hal ini membutuhkan kolaborasi yang lebih erat antara insinyur, peneliti, akademisi, regulator, dan perwakilan industri. Tim yang beragam seperti ini lebih mungkin berkontribusi pada kerangka regulasi AI yang bermakna di industri.

DAFTAR PUSTAKA

- Alvarez-Risco, A., & Del-Aguila-Arcentales, S. (2021). A note on changing regulation in international business: the World Intellectual Property Organization (WIPO) and artificial intelligence. In *The multiple dimensions of institutional complexity in international business research* (pp. 363-371). Emerald Publishing Limited.
- Anderljung, M., Barnhart, J., Korinek, A., Leung, J., O'Keefe, C., Whittlestone, J., ... & Wolf, K. (2023). Frontier AI regulation: Managing emerging risks to public safety. *arXiv preprint arXiv:2307.03718*.
- Beckers, R., Kwade, Z., & Zanca, F. (2021). The EU medical device regulation: Implications for artificial intelligence-based medical device software in medical physics. *Physica Medica*, 83, 1-8.
- Belli, L., Curzi, Y., & Gaspar, W. B. (2023). AI regulation in Brazil: Advancements, flows, and need to learn from the data protection experience. *Computer Law & Security Review*, 48, 105767.
- Bringas Colmenarejo, A., Nannini, L., Rieger, A., Scott, K. M., Zhao, X., Patro, G. K., ... & Kinder-Kurlanda, K. (2022, July). Fairness in agreement with European values: An interdisciplinary perspective on AI regulation. In *Proceedings of the 2022 AAAI/ACM Conference on AI, Ethics, and Society* (pp. 107-118).
- Chae, Y. (2020). US AI regulation guide: legislative overview and practical considerations. *The Journal of Robotics, Artificial Intelligence & Law*, 3.
- Chatterjee, S. (2020). Impact of AI regulation on intention to use robots: From citizens and government perspective. *International Journal of Intelligent Unmanned Systems*, 8(2), 97-114.
- Chatterjee, S., & NS, S. (2023). Impact of AI regulation and governance on online personal data sharing: from sociolegal, technology and policy perspective. *Journal of Science and Technology Policy Management*, 14(1), 157-180.
- Comunale, M., & Manera, A. (2024). The economic impacts and the regulation of AI: A review of the academic literature and policy actions.
- de Laat, P. B. (2021). Companies committed to responsible AI: From principles towards implementation and regulation?. *Philosophy & technology*, 34(4), 1135-1193.
- Durante, M., & Floridi, L. (2022). A legal principles-based framework for AI liability regulation. In *The 2021 yearbook of the digital ethics lab* (pp. 93-112). Cham: Springer International Publishing.
- Ebers, M. (2025). Truly risk-based regulation of artificial intelligence how to implement the EU's AI Act. *European Journal of Risk Regulation*, 16(2), 684-703.
- Fessenko, D. S., & Jasperse, A. (2025). Ethics at the heart of AI regulation. *AI and Ethics*, 5(3), 3387-3398.
- Finocchiaro, G. (2024). The regulation of artificial intelligence. *Ai & Society*, 39(4), 1961-1968.

- Geist, M. A. (2021). AI and international regulation. *Artificial Intelligence and the Law in Canada* (Toronto: LexisNexis Canada, 2021).
- Glauner, P. (2022). An assessment of the ai regulation proposed by the european commission. In *The future circle of healthcare: AI, 3D printing, longevity, ethics, and uncertainty mitigation* (pp. 119-127). Cham: Springer International Publishing.
- Gstrein, O. J. (2022). European AI regulation: Brussels effect versus human dignity?. *Zeitschrift für Europarechtliche Studien (ZEuS)*, 2022(4), 755-772.
- Guan, L., Li, S., & Gu, M. M. (2024). AI in informal digital English learning: A meta-analysis of its effectiveness on proficiency, motivation, and self-regulation. *Computers and Education: Artificial Intelligence*, 7, 100323.
- Hacker, P. (2023). AI regulation in Europe: from the AI act to future regulatory challenges. *arXiv preprint arXiv:2310.04072*.
- Hoffmann, C. H., & Hahn, B. (2020). Decentered ethics in the machine era and guidance for AI regulation. *AI & society*, 35(3), 635-644.
- Hwang, T. J., Kesselheim, A. S., & Vokinger, K. N. (2019). Lifecycle regulation of artificial intelligence– and machine learning–based software devices in medicine. *Jama*, 322(23), 2285-2286.
- Lancieri, F., Edelson, L., & Bechtold, S. (2025). AI regulation: Competition, arbitrage and regulatory capture. *Theoretical Inquiries in Law*, 26(1), 239-262.
- Li, J., Cai, X., & Cheng, L. (2023). Legal regulation of generative AI: a multidimensional construction. *International Journal of Legal Discourse*, 8(2), 365-388.
- Lim, E., Park, H., & Kim, B. (2022, May). Review of the Validity and Rationality of Artificial Intelligence Regulation: Application of the EU's AI Regulation Bill to Accidents Caused by Artificial Intelligence. In *The International FLAIRS Conference Proceedings* (Vol. 35).
- Lucaj, L., Van Der Smagt, P., & Benbouzid, D. (2023, June). Ai regulation is (not) all you need. In *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency* (pp. 1267-1279).
- Maas, M. M. (2022). Aligning AI regulation to sociotechnical change. In *The Oxford Handbook of AI Governance*.
- Meszaros, J., Minari, J., & Huys, I. (2022). The future regulation of artificial intelligence systems in healthcare services and medical research in the European Union. *Frontiers in Genetics*, 13, 927721.
- Middleton, S. E., Letouzé, E., Hossaini, A., & Chapman, A. (2022). Trust, regulation, and human-in-the-loop AI: within the European region. *Communications of the ACM*, 65(4), 64-68.
- Mohebbi, A. (2025). Enabling learner independence and self-regulation in language education using AI tools: a systematic review. *Cogent Education*, 12(1), 2433814.

- Mökander, J., Axente, M., Casolari, F., & Floridi, L. (2022). Conformity assessments and post-market monitoring: a guide to the role of auditing in the proposed European AI regulation. *Minds and Machines*, 32(2), 241-268.
- Molenaar, I. (2022). The concept of hybrid human-AI regulation: Exemplifying how to support young learners' self-regulated learning. *Computers and Education: Artificial Intelligence*, 3, 100070.
- Ngo, D., Nguyen, A., Dang, B., & Ngo, H. (2024). Facial expression recognition for examining emotional regulation in synchronous online collaborative learning. *International Journal of Artificial Intelligence in Education*, 34(3), 650-669.
- Nguyen, A., Järvelä, S., Wang, Y., & Rosé, C. P. (2022). Exploring socially shared regulation with an AI deep learning approach using multimodal data. In *Proceedings of the 16th International Conference of the Learning Sciences-ICLS 2022*, pp. 527-534. International Society of the Learning Sciences.
- Onitiu, D., Wachter, S., & Mittelstadt, B. (2024). How AI challenges the medical device regulation: patient safety, benefits, and intended uses. *Journal of Law and the Biosciences*, Isae007.
- Papyshev, G. (2025). Governing AI through interaction: situated actions as an informal mechanism for AI regulation. *AI and Ethics*, 5(2), 1109-1120.
- Petit, N., & De Cooman, J. (2021). Models of Law and Regulation for AI. *The routledge social science handbook of AI*, 199-221.
- Qiao, H., & Zhao, A. (2023). Artificial intelligence-based language learning: illuminating the impact on speaking skills and self-regulation in Chinese EFL context. *Frontiers in psychology*, 14, 1255594.
- Ridzuan, N. N., Masri, M., Anshari, M., Fitriyani, N. L., & Syafrudin, M. (2024). AI in the financial sector: The line between innovation, regulation and ethical responsibility. *Information*, 15(8), 432.
- Ruscheimer, H. (2023, February). AI as a challenge for legal regulation—the scope of application of the artificial intelligence act proposal. In *Era Forum* (Vol. 23, No. 3, pp. 361-376). Berlin/Heidelberg: Springer Berlin Heidelberg.
- Sartor, G., & Lagioia, F. (2020). The impact of the General Data Protection Regulation (GDPR) on artificial intelligence.
- Seizov, O., & Wulf, A. J. (2020). Artificial intelligence and transparency: a blueprint for improving the regulation of AI applications in the EU. *European Business Law Review*, 31(4).
- Smuha, N. A. (2021). From a 'race to AI' to a 'race to AI regulation': regulatory competition for artificial intelligence. *Law, Innovation and Technology*, 13(1), 57-84.
- Sripada, N. R., Fisher, D. G., & Morris, A. J. (1987, July). AI application for process regulation and servo control. In *IEE Proceedings D (Control Theory and Applications)* (Vol. 134, No. 4, pp. 251-259). IEE.

- Ufert, F. (2020). AI regulation through the lens of fundamental rights: How well does the GDPR address the challenges posed by AI?. *European Papers-A Journal on Law and Integration*, 2020(2), 1087-1097.
- Ulnicane, I. (2022). Artificial Intelligence in the European Union: Policy, ethics and regulation. In *The Routledge handbook of European integrations*. Taylor & Francis.
- Van Kolschooten, H. (2022). EU regulation of artificial intelligence: Challenges for patients' rights. *Common market law review*, 59(1).
- Viljanen, M., & Parviainen, H. (2022). AI applications and regulation: Mapping the regulatory strata. *Frontiers in Computer Science*, 3, 779957.
- Warraich, H. J., Tazbaz, T., & Califf, R. M. (2025). FDA perspective on the regulation of artificial intelligence in health care and biomedicine. *Jama*, 333(3), 241-247.
- Weaver, J. F. (2018). Regulation of artificial intelligence in the United States. In *Research handbook on the law of artificial intelligence* (pp. 155-212). Edward Elgar Publishing.
- Zhao, J. (2024). Promoting more accountable AI in the boardroom through smart regulation. *Computer Law & Security Review*, 52, 105939.

REGULASI APLIKASI AI (Artificial Intelligence)

Dr. Ir. Agus Wibowo, M.Kom, M.Si, MM

BIO DATA PENULIS



Penulis memiliki berbagai disiplin ilmu yang diperoleh dari Universitas Diponegoro (UNDIP) Semarang. dan dari Universitas Kristen Satya Wacana (UKSW) Salatiga. Disiplin ilmu itu antara lain teknik elektro, komputer, manajemen dan ilmu sosiologi. Penulis memiliki pengalaman kerja pada industri elektronik dan sertifikasi keahlian dalam bidang Jaringan Internet, Telekomunikasi, Artificial Intelligence, Internet Of Things (IoT), Augmented Reality (AR), Technopreneurship, Internet Marketing dan bidang pengolahan dan analisa data (komputer statistik).

Penulis adalah pendiri dari Universitas Sains dan Teknologi Komputer (Universitas STEKOM) dan juga seorang dosen yang memiliki Jabatan Fungsional Akademik Lektor Kepala (Associate Professor) yang telah menghasilkan puluhan Buku Ajar ber ISBN, HAKI dari beberapa karya cipta dan Hak Paten pada produk IPTEK. Sejak tahun 2023 penulis tercatat sebagai Dosen luar biasa di Fakultas Ekonomi & Bisnis (FEB) Universitas Diponegoro Semarang. Penulis juga terlibat dalam berbagai organisasi profesi dan industri yang terkait dengan dunia usaha dan industri, khususnya dalam pengembangan sumber daya manusia yang unggul untuk memenuhi kebutuhan dunia kerja secara nyata.



YAYASAN PRIMA AGUS TEKNIK

PENERBIT :

YAYASAN PRIMA AGUS TEKNIK
Jl. Majapahit No. 605 Semarang
Telp. (024) 6723456. Fax. 024-6710144
Email : penerbit_ypat@stekom.ac.id

ISBN 978-634-7227-46-1 (PDF)



9

786347

227461