

Transformasi Layanan Publik Berbasis AI:

Peluang, Risiko, dan Desain Kebijakan
yang Berkeadilan



Oleh Rudy C. Tarumingkeng

Rudy C Tarumingkeng: Transformasi Layanan Publik Berbasis AI -
Peluang, Risiko, dan Desain Kebijakan yang Berkeadilan

Oleh:

[Prof Ir Rudy C Tarumingkeng, PhD](#)

Professor of Management NUP: 9903252922

Rektor, Universitas Cenderawasih, Papua (1978-1988, dan
Rektor, Kampus AGRO Manokwari sekarang Universitas Papua Manokwari)

Coordinator, CIDA/DIKTI SFU Burnaby BC Canada 1988-1991

Rektor, Universitas Kristen Krida Wacana, Jakarta (1991-2000)

Ketua Dewan Guru Besar, IPB-University, Bogor (2005-2006)

AI - Data Analyst, dan Ketua Senat Akademik, IBM-ASMI, Jakarta 2024-

© RudyCT Academic Series

rudyc75@gmail.com

13 Februari 2026

Bogor, Indonesia

TRANSFORMASI LAYANAN PUBLIK BERBASIS AI: PELUANG, RISIKO, DAN DESAIN KEBIJAKAN YANG BERKEADILAN

Prolog: Warga, Antrian, dan “Negara yang Terasa Jauh”

Bayangkan seorang warga—sebut saja Ibu Sari—yang tinggal di pinggiran kota. Ia harus mengurus perpanjangan dokumen kependudukan, mendaftarkan anak ke sekolah negeri, dan memastikan akses bantuan sosial karena pendapatan rumah tangga sedang turun. Ia membuka portal layanan pemerintah, tetapi bahasa kebijakan sulit dipahami, formulir panjang, dan setiap kesalahan kecil membuat proses berulang. Ia datang ke kantor layanan: antrian panjang, petugas terbatas, informasi berbeda antar-loket, dan “sistem sedang gangguan” menjadi kalimat yang terlalu sering terdengar. Dalam pengalaman semacam ini, negara hadir—tetapi terasa jauh: hadir sebagai prosedur, bukan sebagai pelayanan.

Di titik inilah gagasan transformasi layanan publik muncul: bagaimana layanan publik menjadi *lebih cepat, lebih konsisten, lebih mudah dipahami*, sekaligus *lebih adil* bagi warga dengan latar belakang berbeda. Perubahan ini selama dua dekade terakhir didorong oleh e-government dan digitalisasi; namun dalam beberapa tahun terakhir, *Artificial Intelligence (AI)*—termasuk pembelajaran mesin, analitik prediktif, pemrosesan bahasa alami, dan generative AI—membuka lompatan baru: layanan yang tidak sekadar “online”, melainkan “cerdas”, adaptif, dan proaktif.

Namun, di balik janji efisiensi, AI membawa risiko yang tidak kecil: bias, pelanggaran privasi, keputusan yang tidak transparan, otomatisasi yang mengurangi martabat manusia, sampai potensi ketidakadilan sistemik yang sulit dideteksi. Maka pertanyaan kuncinya bukan “apakah pemerintah harus memakai AI?”, melainkan: **bagaimana AI dipakai untuk memperkuat kualitas layanan sekaligus menjamin keadilan, akuntabilitas, dan perlindungan hak warga?**

Esai ini membahas (1) peluang strategis AI bagi layanan publik, (2) spektrum risiko dan dampaknya terhadap keadilan, serta (3) desain kebijakan dan tata kelola (governance) yang berkeadilan—dengan menautkan praktik internasional dan kerangka regulasi yang relevan, serta menerjemahkannya ke konteks Indonesia.

1) AI dalam Layanan Publik: Apa yang Berubah Secara Substantif?

1.1. AI sebagai Lapisan Baru dalam SPBE

Dalam kerangka Indonesia, transformasi digital pemerintahan umumnya diletakkan dalam konsep SPBE. Dalam Peraturan Presiden tentang SPBE, SPBE dipahami sebagai penyelenggaraan pemerintahan yang memanfaatkan teknologi informasi dan komunikasi untuk memberikan layanan kepada pengguna SPBE. Di dalam kerangka yang sama, lahir istilah tata kelola, arsitektur, integrasi proses, data, infrastruktur, dan keamanan sebagai fondasi agar layanan publik terintegrasi.

AI—dalam konteks ini—bukan pengganti SPBE, melainkan **lapisan kemampuan** yang dapat:

mengotomatisasi proses administratif (mis. verifikasi dokumen, klasifikasi permohonan),

meningkatkan kualitas keputusan (mis. deteksi anomali, skor risiko),

memperbaiki antarmuka layanan (mis. chatbot, penerjemah otomatis, ringkasan dokumen),

dan membuka layanan proaktif (mis. pengingat berbasis risiko, prioritisasi kasus mendesak).

Dengan kata lain, bila digitalisasi generasi awal memindahkan layanan ke kanal elektronik, maka AI berpotensi mengubah *logika layanan*: dari reaktif menjadi proaktif, dari "satu ukuran untuk semua" menjadi lebih personal, dari "mengejar berkas" menjadi "mengelola kasus".

1.2. Dari "Layanan" ke "Keputusan": Titik Sensitif AI

AI paling aman dipakai ketika membantu *layanan informasional* (mis. menjawab pertanyaan, membantu navigasi prosedur). Tetapi AI menjadi sensitif ketika dipakai untuk **mendukung atau mengambil keputusan**: siapa yang berhak menerima bantuan, siapa yang harus diperiksa pajaknya, siapa yang diprioritaskan dalam layanan kesehatan, atau siapa yang dianggap "berisiko".

Transformasi layanan publik berbasis AI hampir selalu bergerak dari tahap ringan ke tahap berat:

Augmentasi informasi (AI membantu petugas/warga memahami aturan).

Augmentasi proses (AI mempercepat pekerjaan rutin).

Augmentasi keputusan (AI memberi rekomendasi).

Otomatisasi keputusan (AI memicu tindakan/keputusan administratif).

Semakin ke kanan, semakin besar kebutuhan kebijakan, pengawasan, dan perlindungan hak warga. Di sinilah isu "berkeadilan" menjadi pusat, bukan ornamen.

2) Peluang Strategis AI untuk Layanan Publik

Bagian ini tidak sekadar daftar manfaat, tetapi memetakan *mekanisme* bagaimana AI menghasilkan nilai publik, disertai ilustrasi naratif.

2.1. Akses dan Inklusi: Layanan yang “Bisa Diajak Bicara”

Peluang utama AI di sektor publik adalah menurunkan hambatan akses: bahasa, literasi, disabilitas, dan jarak.

Ilustrasi kasus (naratif):

Pak Andi, pedagang kecil, ingin mengurus izin usaha mikro. Ia tidak memahami istilah hukum dan takut salah mengisi formulir. Dalam skenario AI yang baik, portal layanan menyediakan asisten percakapan yang:

menanyakan kebutuhan Pak Andi dalam bahasa sederhana,
memecah proses menjadi langkah kecil,
memberi contoh pengisian,
memeriksa kelengkapan dokumen,
dan menjelaskan konsekuensi pilihan secara jelas.

Bagi warga, ini bukan sekadar “chatbot lucu”, melainkan *penerjemah negara*. Jika dirancang baik, AI dapat meningkatkan inklusi—terutama bagi kelompok yang selama ini tertinggal dalam layanan digital.

Catatan kebijakan: manfaat inklusi ini hanya terjadi jika AI dilatih/dirancang dengan memperhatikan ragam bahasa, konteks lokal, dan aksesibilitas; jika tidak, AI justru memperlebar jurang digital.

2.2. Efisiensi dan Kecepatan: Mengurangi Biaya Koordinasi yang Tak Terlihat

Dalam organisasi publik, keterlambatan layanan sering disebabkan oleh “biaya koordinasi”: berkas berpindah, unit berbeda tidak sinkron, data tidak terintegrasi, dan prosedur tidak konsisten. AI dapat membantu melalui:

triase otomatis permohonan (mengelompokkan kasus sederhana vs kompleks),

deteksi duplikasi,
ekstraksi informasi dari dokumen,
dan otomatisasi proses rutin.

Ilustrasi kasus (naratif):

Sebuah dinas menerima ribuan permohonan layanan per bulan. Sebelum AI, staf memeriksa satu per satu, membuat antrean menumpuk. Dengan AI yang tepat, permohonan masuk diberi label: *lengkap, kurang dokumen, indikasi salah format, memerlukan verifikasi lapangan*. Staf fokus pada kasus kompleks; kasus sederhana diproses cepat. Dampaknya bukan hanya cepat, tetapi juga konsisten—warga tidak bergantung pada “siapa petugasnya hari ini”.

2.3. Kualitas dan Konsistensi: Mengurangi Variasi Layanan Antar Loket

Dalam layanan publik, ketidakadilan sering muncul bukan karena aturan berbeda, tetapi karena **interpretasi berbeda**. AI dapat membantu menyediakan “pengetahuan prosedural yang sama” untuk semua petugas: rekomendasi langkah, daftar cek, rujukan aturan, dan ringkasan kebijakan. Pada level tertentu, AI dapat mengurangi *street-level discretion* yang tidak perlu—tanpa menghilangkan ruang empati manusia untuk kasus khusus.

Di sinilah generative AI dapat dipakai secara aman: **sebagai mesin ringkas dan mesin rujuk**, bukan mesin memutuskan. Tetapi syaratnya jelas: AI harus diberi sumber rujukan resmi (mis. *retrieval* dari dokumen kebijakan) dan ada mekanisme verifikasi.

2.4. Deteksi Kecurangan dan Perlindungan Anggaran: Integritas sebagai Nilai Publik

AI kuat untuk menemukan pola anomali: klaim yang tidak wajar, jaringan transaksi mencurigakan, atau pola penerima bantuan yang tidak

konsisten. Secara teoritis, ini dapat melindungi anggaran dan meningkatkan integritas.

Namun, di sinilah risiko diskriminasi sering muncul: sistem deteksi “risiko” mudah berubah menjadi sistem “profiling” yang menstigma kelompok tertentu. Karena itu, manfaat integritas harus dipasangkan dengan prinsip keadilan (bagian 3).

2.5. Layanan Proaktif: Negara yang Mengingat, Bukan Menunggu

Salah satu lompatan terbesar adalah layanan proaktif: pemerintah *menghubungi warga* ketika syarat terpenuhi atau risiko meningkat—misalnya pengingat jatuh tempo, rekomendasi program kesehatan, atau bantuan administratif otomatis.

Transformasi ini menuntut ekosistem data yang kuat. Kerangka “Satu Data Indonesia” menggarisbawahi pentingnya tata kelola data untuk menghasilkan data yang akurat, mutakhir, terpadu, dan dapat dipertanggungjawabkan. Tanpa integrasi data yang rapi, layanan proaktif berpotensi salah sasaran—dan kesalahan layanan proaktif jauh lebih sensitif karena “negara datang duluan”.

3) Risiko dan Dilema: Ketika Efisiensi Mengancam Keadilan

Jika peluang AI adalah meningkatkan kapasitas negara, maka risikonya adalah **menciptakan ketidakadilan yang lebih cepat, lebih luas, dan lebih sulit dilawan**. Risiko AI bukan sekadar teknis; ia bersifat sosial, politik, dan etis.

3.1. Bias dan Diskriminasi: Ketidakadilan yang Tersembunyi dalam Data

Bias dapat masuk dari:

data historis yang sudah timpang,

definisi target yang keliru (apa yang dianggap “berisiko”),

variabel proksi (mis. alamat/riwayat tertentu menjadi substitusi kelas sosial),

dan ketidakseimbangan representasi.

Dalam layanan publik, bias bukan sekadar “akurasi menurun”—melainkan **hak warga** terganggu. Jika AI digunakan untuk prioritas layanan, audit, atau bantuan sosial, bias bisa berarti:

warga tertentu lebih sering diperiksa,

kelompok tertentu lebih sering ditolak,

atau akses layanan menjadi tidak setara.

3.2. Privasi dan Perlindungan Data: Risiko “Negara Terlalu Tahu”

AI membutuhkan data; semakin “cerdas”, seringkali semakin “lapar data”. Jika tidak dibatasi, AI mendorong:

pengumpulan data berlebihan,

integrasi lintas instansi tanpa pembatasan,

dan penggunaan ulang data di luar tujuan awal (*function creep*).

Di Indonesia, penguatan etika dan tata kelola pemanfaatan AI sudah mulai dinyatakan melalui pedoman etika yang menekankan nilai, panduan pelaksanaan, dan akuntabilitas pemanfaatan AI bagi pelaku dan penyelenggara sistem elektronik, termasuk lingkup publik. ([JDIH Kemkomdigi](#)) Namun, pedoman etika harus diterjemahkan menjadi mekanisme operasional: minimisasi data, tujuan pemrosesan yang jelas, kontrol akses, audit, dan keamanan siber.

3.3. “Black Box” dan Defisit Akuntabilitas

Warga berhak memahami mengapa sebuah keputusan terjadi. AI yang sulit dijelaskan menciptakan defisit akuntabilitas: keputusan terasa final tetapi tidak dapat diperdebatkan.

Di sektor publik, ini berbahaya karena:

keputusan administrasi memengaruhi hak (bantuan, izin, layanan), dan “ketidakjelasan” memperlemah kepercayaan publik.

Jika AI dipakai untuk rekomendasi, publik perlu tahu: apakah ini rekomendasi atau keputusan? siapa penanggung jawab? bagaimana cara menggugat?

3.4. Automation Bias dan Erosi Pertimbangan Profesional

Risiko lain: petugas terlalu percaya pada sistem. Ketika AI memberi skor risiko, petugas cenderung mengikuti tanpa berpikir kritis. Akibatnya, tanggung jawab moral dan profesional menguap: “sistem yang bilang”.

Kebijakan berkeadilan harus menjaga agar AI memperkuat kapasitas manusia, bukan menggantikannya secara buta.

3.5. Keamanan Siber dan Serangan terhadap Model

Model AI dapat diserang: manipulasi data, *prompt injection* pada sistem generatif, atau eksploitasi kelemahan untuk menghasilkan keputusan salah. Dalam layanan publik, serangan semacam ini bukan sekadar kerugian finansial; bisa mengganggu layanan dasar warga.

3.6. Ketimpangan Akses dan “Keadilan Digital”

AI dapat memperlebar ketimpangan bila:

layanan AI hanya optimal bagi warga dengan internet cepat, antarmuka tidak ramah lansia/disabilitas, atau bahasa lokal tidak terakomodasi.

Keadilan bukan hanya tentang “hasil keputusan” tetapi juga “kesetaraan akses untuk berpartisipasi dalam proses layanan”.

4) Apa Itu “Desain Kebijakan yang Berkeadilan”?

“Berkeadilan” tidak cukup berarti “tidak bias”. Dalam layanan publik, keadilan mencakup beberapa dimensi:

Keadilan distributif: apakah manfaat layanan (kecepatan, bantuan, akses) terdistribusi setara?

Keadilan prosedural: apakah prosesnya transparan, dapat dipahami, dan dapat digugat?

Keadilan korektif: apakah ada mekanisme pemulihan ketika sistem salah?

Keadilan pengakuan: apakah sistem menghormati martabat warga dan keragaman konteks sosial?

Keadilan kapabilitas: apakah layanan memperluas kemampuan warga untuk hidup layak (bukan sekadar memproses berkas)?

Kerangka internasional menegaskan hal serupa: prinsip AI yang “trustworthy” menuntut penghormatan pada hak asasi, transparansi, ketangguhan, dan akuntabilitas. ([OECD Legal Instruments](#)) Prinsip-prinsip ini bukan hanya slogan etika; ia harus menjadi desain kebijakan yang mengikat: standar, audit, kewajiban dokumentasi, dan hak warga.

5) Kerangka Kebijakan Berkeadilan: Dari Prinsip ke Instrumen

Bagian ini adalah inti desain kebijakan: bagaimana menerjemahkan nilai ke perangkat operasional.

5.1. Prinsip Normatif sebagai Kompas

Tiga referensi global yang kuat untuk kompas kebijakan:

Prinsip AI OECD menekankan manfaat bagi manusia dan planet, penghormatan HAM dan rule of law, transparansi, ketangguhan/keamanan, dan akuntabilitas. ([OECD Legal Instruments](#))

Rekomendasi Etika AI UNESCO menekankan HAM, martabat, keadilan, transparansi, dan pentingnya *human oversight* sepanjang siklus hidup AI. ([UNESCO](#))

NIST AI Risk Management Framework memberi pendekatan manajemen risiko yang dapat dipakai lintas sektor untuk meningkatkan trustworthiness dan mengelola risiko pada individu, organisasi, dan masyarakat. ([NIST](#))

Indonesia sendiri mulai mengarahkan etika AI sebagai pedoman bagi pelaku usaha dan penyelenggara sistem elektronik lingkup publik dan privat. ([JDIH Kemkomdigi](#)) Ini penting sebagai “landasan nilai”, tetapi kebijakan berkeadilan membutuhkan perangkat yang lebih konkret.

5.2. Klasifikasi Risiko dan “Proportionality”

Tidak semua penggunaan AI sama. Menggunakan AI untuk menjawab FAQ berbeda risikonya dengan AI untuk menentukan kelayakan bantuan sosial. Uni Eropa memilih pendekatan **risk-based regulation** melalui AI Act, dengan timeline penerapan bertahap dan pembedaan kewajiban berdasarkan tingkat risiko. ([Digital Strategy](#))

Bagi kebijakan layanan publik, ini dapat diterjemahkan menjadi:

Risiko rendah: chatbot informasional, pencarian dokumen, ringkasan kebijakan (dengan batasan).

Risiko menengah: rekomendasi triase layanan, analitik untuk prioritas kasus.

Risiko tinggi: sistem yang memengaruhi hak, akses bantuan, pengawasan warga, atau sanksi administratif.

Risiko terlarang: praktik yang melanggar martabat/HAM (misalnya “skor sosial” gaya tertentu), kecuali ada dasar hukum yang sangat ketat dan uji proporsionalitas yang kuat.

Intinya: **semakin besar dampak terhadap hak warga, semakin ketat kewajiban tata kelolanya.**

5.3. Algorithmic Impact Assessment: Wajib untuk Sistem Berdampak Tinggi

Praktik baik yang bisa diadaptasi adalah **Algorithmic Impact Assessment (AIA)** yang digunakan pemerintah Kanada: sebuah kuesioner penilaian risiko yang menentukan level dampak dan mewajibkan mitigasi sesuai skor. ([Canada](#))

Dalam konteks Indonesia, AIA versi nasional untuk layanan publik dapat mewajibkan:

definisi tujuan layanan (problem framing),
identifikasi pihak terdampak,
analisis potensi diskriminasi,
penilaian privasi dan keamanan,
pilihan desain “human-in-the-loop”,
rencana pemantauan dan mekanisme banding.

Tanpa AIA, transformasi AI sering menjadi proyek teknologi semata; padahal ini proyek kebijakan publik yang menyentuh hak.

5.4. Transparansi Publik: “Daftar AI Pemerintah” yang Bisa Diakses Warga

Salah satu inovasi tata kelola adalah **Algorithmic Transparency Recording Standard (ATRS)** di Britania Raya yang mendorong lembaga publik mempublikasikan informasi tentang alat algoritmik yang digunakan dan alasan penggunaannya. ([GOV.UK](#))

Adaptasi yang relevan: **Registri AI Layanan Publik** yang memuat:
nama sistem,

tujuan dan ruang lingkup,
data yang dipakai,
level risiko,
evaluasi dampak,
cara warga mengajukan keberatan,
dan penanggung jawab institusional.

Transparansi bukan sekadar “buka kode”; transparansi adalah **membuka pengetahuan yang cukup** agar publik bisa mengawasi.

5.5. Manajemen Risiko Sepanjang Siklus Hidup

AI bukan produk sekali jadi. Model berubah karena data baru, kebijakan baru, atau perubahan perilaku publik. Maka tata kelola harus “lifecycle-based”. NIST AI RMF menekankan pengelolaan risiko untuk meningkatkan trustworthiness sepanjang desain, pengembangan, penggunaan, dan evaluasi sistem AI. ([NIST](#))

Instrumen operasional yang lazim:

audit berkala,
pengujian bias dan performa,
pemantauan *drift*,
incident reporting,
dan pembekuan sistem bila melampaui ambang risiko.

5.6. Hak Warga: Penjelasan, Keberatan, dan Pemulihan

Desain berkeadilan harus menjamin minimal:

hak untuk diberi tahu ketika AI digunakan dalam proses layanan,
hak atas penjelasan yang bermakna (bukan jargon teknis),

hak untuk menolak keputusan otomatis tertentu (terutama yang berdampak tinggi),

hak banding kepada petugas manusia,

mekanisme pemulihan (perbaikan data, koreksi hasil, kompensasi bila perlu).

Tanpa mekanisme banding, AI dapat menjadi “otoritas tak terlihat” yang sulit dilawan.

6) Konteks Indonesia: Fondasi Strategis dan Ruang Kebijakan

6.1. SPBE dan Arsitektur Integrasi

Kerangka SPBE menekankan arsitektur yang mengintegrasikan proses bisnis, data dan informasi, infrastruktur, aplikasi, serta keamanan untuk menghasilkan layanan terintegrasi. Bagi AI, ini penting: AI yang berjalan di atas sistem terfragmentasi biasanya memperbesar kesalahan, karena sumber data tidak sinkron.

6.2. Satu Data: Prasyarat Keadilan Berbasis Data

Perpres tentang Satu Data Indonesia menekankan tata kelola dan integrasi data lintas instansi, yang relevan untuk memastikan data layanan publik konsisten dan dapat dipertanggungjawabkan.

Keadilan AI hampir selalu berujung pada kualitas data: jika data kependudukan, ekonomi, atau layanan sosial tidak rapi, maka AI akan “mengeraskan” ketidakrapihan itu menjadi keputusan cepat.

6.3. Arah Strategis AI Nasional dan Etika Pancasila

Sekretariat Kabinet Republik Indonesia memuat bahwa rancangan strategi nasional AI diarahkan selaras dengan visi Indonesia Emas 2045, dengan misi antara lain mewujudkan AI yang beretika sesuai nilai Pancasila, menyiapkan talenta, membangun ekosistem data/infrastruktur, dan mendorong kolaborasi riset untuk reformasi birokrasi dan industri.

([Sekretariat Kabinet Republik Indonesia](#))

Ini memberi legitimasi strategis: AI bukan sekadar proyek TI, melainkan bagian dari reformasi birokrasi dan pembangunan kapasitas negara.

6.4. Pedoman Etika AI dan Langkah Konsultasi Publik

Kementerian Komunikasi dan Digital melalui Surat Edaran tentang Etika Kecerdasan Artifisial memberi pedoman nilai, panduan umum, dan akuntabilitas dalam pemanfaatan/pengembangan AI. ([JDIH Kemkomdigi](#)) Selain itu, terdapat pengumuman konsultasi publik terkait Buku Putih peta jalan AI nasional dan konsep pedoman etika AI. ([Kementerian Komunikasi dan Digital](#))

Secara kebijakan, ini peluang penting untuk membangun legitimasi sosial: AI di layanan publik membutuhkan *social license*, bukan hanya *technical license*.

7) Rancangan Kebijakan Berkeadilan: Paket Desain yang Realistis

Di bagian ini saya susun paket kebijakan yang dapat dianggap “desain institusional”, bukan sekadar daftar imbauan.

7.1. Portofolio Use-Case dan Peta Risiko Nasional

Langkah pertama: pemerintah perlu menginventarisasi use-case AI layanan publik dan memetakannya ke tier risiko. Ini seperti manajemen portofolio kebijakan: mana yang cepat memberi manfaat (quick wins) tanpa risiko hak, dan mana yang harus ditunda sampai perangkat pengawasan matang.

Prinsip:

mulai dari layanan informasional dan proses administratif rendah risiko, lanjut ke rekomendasi,

batasi otomatisasi keputusan untuk area yang menyentuh hak warga kecuali ada dasar hukum dan pengawasan kuat.

7.2. Standar Pengadaan dan Anti Vendor Lock-in

Salah satu sumber risiko adalah pengadaan “black box” dari vendor yang sulit diaudit. Kebijakan pengadaan harus mewajibkan:

dokumentasi model dan data (model card, data sheet),

hak audit pemerintah,

standar interoperabilitas,

rencana keluar (exit plan),

dan pengujian independen untuk sistem berdampak tinggi.

Tanpa ini, AI layanan publik bisa menjadi ketergantungan jangka panjang.

7.3. AIA + DPIA sebagai Paket Wajib Sistem High-Impact

Sistem high-impact harus melewati:

AIA (dampak algoritmik), ([Canada](#))

penilaian dampak perlindungan data (DPIA),

dan uji keadilan (fairness testing) yang dipublikasikan ringkasannya.

Ini menjembatani etika (nilai) menjadi administrasi publik (prosedur).

7.4. Registri Transparansi + Komunikasi Publik yang Sederhana

Mengadopsi semangat ATRS: lembaga publik wajib membuat catatan transparansi alat algoritmik yang digunakan. ([GOV.UK](#))

Bentuk komunikasinya harus ramah warga: ringkas, bahasa sederhana, ada kanal tanya-jawab, dan mekanisme keberatan.

7.5. Desain Human Oversight yang Nyata

“Human-in-the-loop” bukan slogan. Kebijakan perlu mendefinisikan:

kapan petugas boleh menolak rekomendasi AI,

kapan keputusan harus ditinjau manusia,

standar pelatihan petugas,

dan indikator *automation bias* (mis. tingkat pembalikan keputusan).

UNESCO menekankan pentingnya pengawasan manusia sebagai prinsip etika. ([UNESCO](#)) Dalam praktik layanan publik, ini berarti memperkuat kapasitas petugas—bukan sekadar memasang mesin.

7.6. Mekanisme Banding dan Ombudsman Algoritmik

Untuk layanan berdampak tinggi, perlu kanal banding khusus:

warga dapat meminta peninjauan manual,

ada SLA penanganan banding,

dan audit kasus banding digunakan untuk memperbaiki model.

Pilihan institusionalnya bisa beragam: unit internal (inspektorat), badan pengawas lintas instansi, atau memperkuat fungsi ombudsman/komisi terkait. Yang penting: **warga tidak dibiarkan sendirian melawan sistem yang tidak terlihat.**

7.7. Monitoring, Audit, dan Incident Reporting

Mengacu pada pendekatan manajemen risiko seperti NIST AI RMF, tata kelola harus mencakup monitoring berkelanjutan dan pelaporan insiden. ([NIST](#))

Misalnya:

insiden bias (kelompok tertentu dirugikan),

kebocoran data,

kegagalan sistem,

atau “halusinasi” generative AI yang menyesatkan.

Setiap insiden harus punya prosedur: investigasi, perbaikan, dan komunikasi publik proporsional.

8) Ilustrasi Rangkaian Kasus: Dari Chatbot Sampai Keputusan Bantuan

Agar terlihat perbedaan risiko, saya susun tiga ilustrasi naratif bertingkat.

Kasus A (Risiko Rendah): Asisten Informasi Izin Usaha

Sebuah pemerintah daerah meluncurkan asisten AI untuk membantu izin usaha. AI hanya:

- menjelaskan prosedur,
- memberi daftar dokumen,
- membantu mengisi formulir,
- dan mengarahkan kanal resmi.

Kebijakan minimal yang wajib:

- transparansi bahwa ini AI,
- sumber informasi resmi (agar tidak "mengarang"),
- log percakapan dengan perlindungan data,
- dan tombol eskalasi ke petugas manusia.

Kasus B (Risiko Menengah): Triase Permohonan Bantuan dengan Rekomendasi

AI memberi rekomendasi prioritas verifikasi lapangan berdasarkan indikator administratif (kelengkapan, konsistensi data, pola anomali). Keputusan final tetap di petugas.

Kebijakan wajib:

- AIA ringkas,
- pengujian bias (apakah kelompok tertentu selalu "dicurigai"),
- prosedur petugas untuk membatalkan rekomendasi,
- dan audit berkala.

Kasus C (Risiko Tinggi): Penetapan Kelayakan Bantuan Secara Otomatis

AI menentukan “layak/tidak layak” bantuan. Ini menyentuh hak dasar warga.

Kebijakan ketat (high-impact):

dasar hukum jelas,

AIA + DPIA penuh,

hak banding dan peninjauan manusia,

transparansi kriteria,

registri publik,

pengawasan independen,

serta pembatasan variabel (hindari proksi diskriminatif).

Di titik ini, kebijakan berkeadilan menentukan apakah AI menjadi alat perlindungan sosial atau alat eksklusif sosial.

9) Roadmap Implementasi: Reformasi Manajemen, Bukan Sekadar Teknologi

Transformasi layanan publik berbasis AI adalah agenda manajemen perubahan. Tiga aspek yang sering menentukan keberhasilan:

Kapasitas institusi

Talenta data, analisis kebijakan, auditor algoritmik, dan pejabat pengadaan harus ditingkatkan. Ini selaras dengan misi penyiapan talenta dalam arah strategis AI Indonesia. ([Sekretariat Kabinet Republik Indonesia](#))

Arsitektur dan data

AI berkualitas tumbuh di atas data yang tertib—sejalan dengan SPBE dan Satu Data.

Legitimasi publik

Konsultasi publik dan pedoman etika harus menjadi proses membangun kepercayaan, bukan formalitas. ([Kementerian Komunikasi dan Digital](#))

10) Penutup: AI sebagai Ujian Etika Negara Modern

AI dapat membuat negara terasa dekat: layanan cepat, jelas, dan responsif. Tetapi AI juga dapat membuat negara terasa menakutkan: keputusan otomatis yang tidak bisa dipahami, penilaian risiko yang memprofilkan warga, dan pengumpulan data tanpa batas.

Karena itu, “transformasi layanan publik berbasis AI” pada dasarnya adalah **ujian etika dan tata kelola negara modern**. Negara yang berkeadilan bukan negara yang paling canggih algoritmanya, melainkan negara yang:

memakai AI untuk memperluas akses layanan,
membatasi AI ketika menyentuh martabat dan hak warga,
memastikan transparansi dan akuntabilitas,
menyediakan mekanisme banding dan pemulihan,
serta menjaga agar manusia tetap menjadi pusat keputusan yang bermoral.

Dalam bahasa kebijakan: **AI harus menjadi alat untuk memperkuat keadilan layanan publik, bukan mesin yang mempercepat ketidakadilan.**

Berikut **Glosarium** (istilah kunci) dan **Referensi** terpilih untuk topik “*Transformasi Layanan Publik Berbasis AI: Peluang, Risiko, dan Desain Kebijakan yang Berkeadilan.*”

Glosarium

AI (Artificial Intelligence) / Kecerdasan Artifisial: Sistem komputasi yang mampu melakukan tugas yang biasanya membutuhkan kecerdasan manusia (mis. klasifikasi, prediksi, rekomendasi, dialog).

Machine Learning (ML): Cabang AI yang “belajar” dari data untuk membuat prediksi/keputusan tanpa aturan yang diprogram secara eksplisit.

Deep Learning: Subset ML yang menggunakan jaringan saraf berlapis (neural networks) untuk menangkap pola kompleks; efektif untuk teks, suara, citra.

Generative AI: AI yang menghasilkan konten baru (teks, gambar, audio, kode) berdasarkan pola dari data pelatihan.

LLM (Large Language Model): Model generatif berbasis teks yang dilatih pada korpus besar untuk memahami dan menghasilkan bahasa alami.

NLP (Natural Language Processing): Teknik AI untuk memproses bahasa manusia (klasifikasi dokumen, chatbot layanan publik, ringkasan kebijakan).

Computer Vision: AI untuk memahami citra/video (mis. deteksi kepadatan antrean, verifikasi dokumen).

Automated Decision System (ADS): Sistem yang mengotomasi keputusan/kelayakan (mis. bantuan sosial, prioritas layanan), baik sepenuhnya otomatis maupun semi-otomatis.

Decision Support System (DSS): Sistem pendukung keputusan; manusia tetap memutuskan, AI memberi rekomendasi/insight (mis. triase layanan).

Risk-based Approach (Pendekatan berbasis risiko): Kebijakan/standar yang mensyaratkan kontrol lebih ketat pada penggunaan AI berisiko tinggi (mis. menyangkut hak warga).

High-stakes / High-impact Use Case: Use case AI yang berdampak besar terhadap hak, akses layanan, keselamatan, dan martabat manusia (mis. penentuan penerima bansos).

Bias (Bias Algoritmik): Penyimpangan sistematis yang menyebabkan hasil merugikan kelompok tertentu; dapat berasal dari data, desain model, atau konteks implementasi.

Disparate Impact (Dampak timpang): Dampak kebijakan/algoritme yang secara statistik merugikan kelompok tertentu meski tanpa niat diskriminasi. ([Yale Computer Science](#))

Fairness (Keadilan/Fairness dalam AI): Keluarga konsep & metrik untuk mengurangi ketimpangan hasil (mis. kesetaraan error antar kelompok), dengan trade-off yang perlu diputuskan secara normatif.

Explainability (Keterjelasan alasan): Upaya membuat alasan keluaran model dapat dipahami; penting untuk akuntabilitas layanan publik.

Interpretability (Keterpahaman model): Seberapa mudah mekanisme model dipahami (berbeda dari explainability yang sering “pasca-hoc”).

Transparency (Transparansi): Keterbukaan tentang *apa* sistemnya, *untuk apa*, *data apa*, *siapa vendor*, *bagaimana diuji*, dan *bagaimana warga mengajukan keberatan*.

Accountability (Akuntabilitas): Kejelasan penanggung jawab (pejabat/instansi) atas keputusan berbasis AI dan konsekuensinya—bukan “menyalahkan mesin”.

Algorithmic Audit (Audit Algoritmik): Pemeriksaan sistematis atas model/data/proses untuk mendeteksi bias, kesalahan, dan risiko; termasuk audit internal end-to-end. ([ACM Digital Library](#))

Model Card: Dokumen standar yang menjelaskan tujuan, konteks, performa, keterbatasan, dan isu etis sebuah model. ([arXiv](#))

Datasheet for Datasets: Dokumentasi standar dataset: motivasi, komposisi, proses pengumpulan, risiko, dan batasan pemakaian. ([arXiv](#))

AIA (Algorithmic Impact Assessment): Instrumen penilaian dampak/risiko sistem keputusan otomatis untuk menentukan mitigasi dan kewajiban transparansi. ([Canada](#))

DPIA (Data Protection Impact Assessment): Kajian dampak pemrosesan data pribadi berisiko tinggi (privasi & hak subjek data) sebelum sistem dijalankan (praktik umum dalam rezim perlindungan data).

Data Governance (Tata kelola data): Kebijakan, peran, standar, dan kontrol kualitas data (ownership, akses, metadata, lineage).

Data Quality (Kualitas data): Akurasi, kelengkapan, konsistensi, ketepatan waktu, dan keterlacakan data—kritis bagi fairness dan ketepatan keputusan.

Satu Data Indonesia: Kerangka kebijakan integrasi data pemerintah dengan prinsip standar data, metadata, interoperabilitas, dan tata kelola. ([BPK Regulations](#))

SPBE (Sistem Pemerintahan Berbasis Elektronik): Kerangka transformasi digital pemerintahan untuk layanan publik yang efektif, transparan, dan akuntabel. ([BPK Regulations](#))

Interoperability (Interoperabilitas): Kemampuan sistem lintas instansi untuk bertukar dan menggunakan data secara konsisten (API, standar data, identitas).

PSE Lingkup Publik: Penyelenggara sistem elektronik pada instansi negara; wajib memperhatikan aspek keamanan/keandalan sesuai regulasi. ([BPK Regulations](#))

Privacy by Design: Prinsip merancang privasi sejak awal (minimisasi data, kontrol akses, logging, retensi terbatas).

Data Minimization (Minimisasi data): Mengambil data seperlunya untuk tujuan tertentu; mengurangi risiko kebocoran dan penyalahgunaan.

Purpose Limitation (Pembatasan tujuan): Data diproses untuk tujuan yang jelas; perlu kebijakan jika tujuan berubah.

Anonymization vs Pseudonymization: Anonimisasi menghilangkan identitas secara kuat; pseudonimisasi mengganti identitas dengan kode namun masih bisa direkonstruksi dengan kunci.

Cybersecurity (Keamanan siber): Perlindungan sistem dari serangan; relevan untuk layanan publik digital.

ISMS (Information Security Management System): Sistem manajemen keamanan informasi; salah satu rujukan global adalah ISO/IEC 27001. ([ISO](#))

Human-in-the-Loop (HITL): Manusia ikut menilai/menyetujui keputusan AI (khususnya high-stakes).

Human-on-the-Loop: Manusia mengawasi sistem yang berjalan, menetapkan ambang, dan intervensi saat anomali.

Recourse & Contestability (Keberatan dan pemulihan): Mekanisme warga untuk menantang keputusan, meminta peninjauan, dan memperoleh koreksi/kompensasi.

Monitoring & Drift: Pemantauan kinerja dan perubahan pola data (drift) yang bisa membuat model “menurun” dan bias baru muncul.

Vendor Lock-in: Ketergantungan pada vendor/produk tertentu sehingga sulit pindah; berimplikasi pada biaya, kedaulatan data, dan akuntabilitas.

Referensi

A. Regulasi & Kebijakan Indonesia (fondasi tata kelola layanan publik digital)

Republik Indonesia. **Undang-Undang No. 25 Tahun 2009 tentang Pelayanan Publik.** ([Mahkamah Agung RI](#))

Republik Indonesia. **Peraturan Presiden No. 95 Tahun 2018 tentang Sistem Pemerintahan Berbasis Elektronik (SPBE).** ([BPK Regulations](#))

Republik Indonesia. **Peraturan Presiden No. 39 Tahun 2019 tentang Satu Data Indonesia.** ([BPK Regulations](#))

Republik Indonesia. **Undang-Undang No. 27 Tahun 2022 tentang Pelindungan Data Pribadi.** ([BPK Regulations](#))

Republik Indonesia. **Peraturan Pemerintah No. 71 Tahun 2019 tentang Penyelenggaraan Sistem dan Transaksi Elektronik.** ([BPK Regulations](#))

Republik Indonesia. **Undang-Undang No. 14 Tahun 2008 tentang Keterbukaan Informasi Publik.** ([JDIH ESDM](#))

Kementerian Komunikasi dan Digital (Komdigi). **Surat Edaran Menteri Komunikasi dan Informatika No. 9 Tahun 2023 tentang Etika Kecerdasan Artifisial.** ([JDIH Kemkomdigi](#))

Kementerian Komunikasi dan Digital (Komdigi). **Konsultasi publik Buku Putih Peta Jalan Kecerdasan Artifisial Nasional & konsep pedoman etika** (pengumuman/siaran pers). ([Kementerian Komunikasi dan Digital](#))

B. Kerangka Etika, Standar, dan Tata Kelola Global (praktik baik untuk desain kebijakan berkeadilan)

OECD. **Recommendation of the Council on Artificial Intelligence (OECD/LEGAL/0449)** (2019). ([OECD Legal Instruments](#))

UNESCO. **Recommendation on the Ethics of Artificial Intelligence** (2021). ([UNESCO Digital Library](#))

NIST. **Artificial Intelligence Risk Management Framework (AI RMF 1.0)** (2023). ([NIST Publications](#))

European Commission. **AI Act enters into force** (berita resmi, 1 Aug 2024) dan laman kebijakan AI Act (pembaruan 27 Jan 2026). ([European Commission](#))

United Kingdom. **Algorithmic Transparency Recording Standard (ATRS): Guidance for public sector bodies**. ([GOV.UK](#))

Canada. **Algorithmic Impact Assessment tool** dan **Guide on the Scope of the Directive on Automated Decision-Making**. ([Canada](#))

International Organization for Standardization (ISO). **ISO/IEC 27001 (ISMS) overview**. ([ISO](#))

C. Literatur Akademik Kunci (bias, akuntabilitas, dokumentasi, audit)

Barocas, S., & Selbst, A. D. **Big Data's Disparate Impact**. *California Law Review* (2016). ([Yale Computer Science](#))

Kroll, J. A., dkk. **Accountable Algorithms**. *University of Pennsylvania Law Review* (2017). ([scholarship.law.upenn.edu](#))

Mitchell, M., dkk. **Model Cards for Model Reporting**. (arXiv 2018; ACM 2019). ([arXiv](#))

Gebru, T., dkk. **Datasheets for Datasets**. (arXiv 2018). ([arXiv](#))

Raji, I. D., dkk. **Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing**. (arXiv/ACM 2020). ([ACM Digital Library](#))

D. Buku Rujukan (narasi dampak sosial—berguna untuk studi kasus layanan publik)

Automating Inequality. Eubanks, V. **Automating Inequality: How High-Tech Tools Profile, Police, and Punish the Poor** (2018). ([Virginia Eubanks](#))

Weapons of Math Destruction. O'Neil, C. **Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy** (2016). ([Google Books](#))

Race After Technology. Benjamin, R. **Race After Technology: Abolitionist Tools for the New Jim Code** (2019). ([ruhabenjamin.com](#))

Copilot for this article - Chatgpt 5.2 Thinking. Access date: 13 Februari 2026. Prompting on Writer's account ([Rudy C Tarumingkeng](#)) <https://chatgpt.com/c/698e8973-6cbc-839b-aa1b-02cf2cd4d442>