

## *Is There a Constitutional Right to a Human Decision in Canada?*

Robert Diab\*

*In the mid-2010s, government interest in and uptake of automated decision-making prompted debate over law and policy on point. At issue were forms of AI that did not give reasons and relied on algorithms of varying degrees of opacity. Canada's Treasury Board Directive on Automated Decision-Making and Guide on the Use of Generative Artificial Intelligence follow the EU's GDPR and AI Act in recognizing rights to a human decision and to an explanation of automated decision-making, thereby limiting the use of AI in high-impact cases. But recent advances in AI challenge the assumptions underlying this approach. Language models can generate reasons for a decision comparable in quality to those of a human — prompting a need to look beyond existing law and policy toward constitutional limits on use of this technology. This article outlines the contours of Charter rights to a human decision and an explanation, which, for the foreseeable future, preclude significant reliance on language models in high-impact administrative cases regardless of their capabilities.*

### **I. Introduction**

In the mid-2010s, government interest in automated decision-making (“ADM”) in administrative law was on the rise in Canada, the European Union, and elsewhere, prompting a

---

\* Professor of Law, Thompson Rivers University.

debate over law and policy on point.<sup>1</sup> The policy debate was shaped in part by the state of the technology at issue. Decisions made in partial or full reliance on algorithms tended to involve software that generated scores or predictions (recidivism, credit worthiness) without providing reasons.<sup>2</sup> The variables, weights, or factors shaping the output fell somewhere along a spectrum of transparency from clear and simple algorithms to machine learning processes that were largely opaque.<sup>3</sup> Law and policy generally accepted a limited role for ADM at the more transparent end of the spectrum and took a skeptical view of ADM involving the latter.

This is reflected Europe's approach to regulating ADM, which has formed an important backdrop to debates in Canada. The EU's *General Data Protection Regulation* (2016) and *Artificial Intelligence Act* (2024) contained the basis for rights to a human decision and a "meaningful" explanation of an automated decision.<sup>4</sup> These cast doubt on AI's possible role in high-impact cases, due in part to the opacity of its underlying processes. Canada imposes similar requirements in its Directive on Automated Decision-Making,<sup>5</sup> and its Guide on the Use of Generative Artificial Intelligence specifically rules out using language models for ADM, again due in part to concerns about algorithmic opacity.<sup>6</sup>

This would seem to settle questions about the legal limits on the use of AI in ADM in Canada. But recent experiments with language models present a new possibility that challenges earlier assumptions about AI unsuitability due to opacity: the possibility of AI that can explain why, though not how, it reached a certain decision. Experiments show that AI drawing on party materials can generate detailed and persuasive reasons to justify a given outcome — and can do so on a level comparable in quality to that of a human.<sup>7</sup> Other developments point to a near future in which AI might soon acquire the ability to hear cases, including oral

---

1 Amy Gouge, "Administrative Law, Artificial Intelligence, and Procedural Rights" (2021) 42 Windsor Rev Legal Soc Issues 17; Jennifer Raso, "AI and Administrative Law" in Florian Martin-Bariteau & Teresa Scassa, eds, *Artificial Intelligence and the Law in Canada* (Toronto: LexisNexis Canada, 2021); Cary Coglianese & David Lehr, "Regulating by Robot: Administrative Decision Making in the Machine-Learning Era" (2017) 105:5 Geo LJ 1147.

2 See e.g. Gouge, *supra* note 1 at 24–5, discussing AI used "to fix credit scores," "predict tax evasion," "detect welfare fraud," and "evaluate risk in criminal offenders"; Raso, *supra* note 1, focuses on risk scores used in the prison security classification context.

3 For discussions of this distinction, see Gouge, *supra* note 1 at 24; Paul Daly, "Artificial Administration: Administrative Law, Administrative Justice and Accountability in the Age of Machines" (2023) 30:2 Australian Journal of Administrative Law & Practice 95; Jesse Beatson, "AI-Supported Adjudicators: Should Artificial Intelligence Have a Role in Tribunal Adjudication?" (2018) 31 Can J Admin L & Prac 307 at 308.

4 Regulation (EU) 2016/679 of the European Parliament and of the Council of 27 April 2016: *on the protection of natural persons with regard to the processing of personal data and on the free movement of such data, and repealing Directive 95/46/EC (General Data Protection Regulation)*, [2016] OJ L 119/1 at 1–88 [GDPR]; Regulation (EU) 2024/1243 of the European Parliament and of the Council of 11 July 2024: *laying down harmonised rules on artificial intelligence (Artificial Intelligence Act) [AI Act]*.

5 Treasury Board of Canada Secretariat, "Directive on Automated Decision-Making" (5 February 2019, updated in June 2025), online: *Government of Canada* <<https://www.tbs-sct.gc.ca/pol/doc-eng.aspx?id=32592>> [<https://perma.cc/RE86-HQBB>] [DADM].

6 Treasury Board of Canada Secretariat, "Guide on the Use of Generative Artificial Intelligence" (September 2023, updated February 2024), online: *Government of Canada* <<https://www.canada.ca/en/government/system/digital-government/digital-government-innovations/responsible-use-ai/guide-use-generative-ai.html>> [Guide].

7 Details are discussed in Part IV below.

evidence, and to replace human adjudicators in a wider range of venues.<sup>8</sup> These possibilities prompt a need to look beyond existing law and policy and to address an issue under-explored in the literature: is there a constitutional right in Canada to a human decision and explanation that would impose an upper limit on the use of AI in ADM, regardless of its capabilities?

This article briefly canvasses the law and policy noted above, along with recent experiments with language models in judgment, to provide context for an outline of relevant rights under the *Canadian Charter of Rights and Freedoms* (“*Charter*”).<sup>9</sup> The paper argues that courts are likely to interpret elements of common law procedural fairness that have been constitutionalized as components of fundamental justice under section 7 of the *Charter* — such as a hearing by an impartial tribunal, notice of the case to meet, and the provision of reasons — as precluding significant reliance on generative AI for ADM, unless and until prevailing understandings of the opacity underlying AI change.

## II. European Law on the Right to a Human Decision and an Explanation

European regulation of automated decision-making provides an important backdrop to Canadian debates about rights to a human decision and an explanation. These issues are addressed principally in the *General Data Protection Regulation* (“GDPR”) <sup>10</sup> and, more recently, the EU’s *Artificial Intelligence Act* (“*AI Act*”),<sup>11</sup> both of which reflect assumptions about transparency, human oversight, and the limits of automation in high-impact decisions.

Article 22 of the GDPR restricts decisions “based solely on automated processing” that produce legal or similarly significant effects.<sup>12</sup> While the Regulation permits exceptions — such as where an EU Member State authorizes automated decisions in law or where a person consents — it requires “suitable safeguards,” including at minimum a right to human intervention, to express one’s point of view, and to contest the decision.<sup>13</sup> Although Article 22 itself does not explicitly guarantee a right to an explanation, Recital 71 indicates that safeguards should also include providing “specific information” about how a decision was reached.

Further articles of the GDPR require that individuals subject to automated decision-making be informed of the fact and be provided with “meaningful information about the logic involved.”<sup>14</sup> Guidance issued by a Working Party convened under the GDPR clarifies that this does not require disclosure of algorithms, but it does require explanations sufficient for individuals to understand the reasons for a decision.<sup>15</sup> In a case concerning automated credit scor-

---

8 See the developments on the auditory front in Changli Tang et al, “SALMONN: Towards Generic Hearing Abilities for Large Language Models” (8 April 2024), online, arXiv: <<https://arxiv.org/abs/2310.13289>> [<https://perma.cc/X6P6-KL5B>]; for developments in visual perception capabilities of generative AI, see Adam Clark Estes, “An AI that sees what you see,” *Vox* (12 December 2024).

9 *Canadian Charter of Rights and Freedoms*, Part I of the *Constitution Act, 1982*, being Schedule B to the *Canada Act 1982* (UK), 1982, c 11 [*Charter*].

10 GDPR, *supra* note 4.

11 *AI Act*, *supra* note 4.

12 GDPR, *supra* note 4, Art 22(1).

13 *Ibid*, Art 22(3).

14 GDPR, *supra* note 4, Arts 13(2)(f), 14(2)(g), 15(1)(h).

15 Article 29 Data Protection Working Party, *Guidelines on Automated Individual Decision-Making and Profiling for the Purposes of Regulation 2016/679* (wp251rev.01) (6 February 2018).

ing, the Court of Justice of the European Union held that transparency requirements may be satisfied through counterfactual explanations — what would have led to a different outcome? — rather than disclosure of technical processes.<sup>16</sup>

The *AI Act* builds on the GDPR framework. While it does not establish an explicit right to a human decision, it classifies AI systems used to make legally significant decisions as “high-risk” and mandates human oversight, interpretability, and the ability to override outputs.<sup>17</sup> The *Act* also contemplates that AI tools may support judicial decision-making but a recital stipulates that this must not replace human judgment.<sup>18</sup> Article 86 further establishes a limited right to an explanation of the role played by AI in adverse decisions.

Taken together, these components of EU law preserve space for the use of AI in decision-making, but condition it on meaningful human involvement and explanations that need not fully penetrate the opacity of complex machine-learning systems.

### III. Canada’s DADM and Guide on Generative AI

Two important policy documents in Canada on ADM follow the path set out in EU law. They form further context for considering how courts might construe rights to a human decision and an explanation. Neither the Treasury Board’s “Directive on Automated Decision-Making” (“DADM”),<sup>19</sup> nor the Board’s “Guide on the Use of Generative Artificial Intelligence”<sup>20</sup> has the force of law, and both apply only to automated decisions by federal institutions and agencies.<sup>21</sup> They can, however, be read as reflective of normative assumptions about the use of AI in judgment that would inform questions about legitimate expectations around AI when addressing questions of procedural fairness at common law and “fundamental justice” under section 7 of the *Charter* (canvassed below). For this purpose, I highlight only a few key features of both documents.

First issued in 2019 and updated in 2025, the DADM contemplates the possible use of AI in the provision of a service to a person external to government to make a decision “in whole or in part,” so long as notice of this is provided,<sup>22</sup> affected persons receive a “meaningful explanation ... of how and why the decision was made,”<sup>23</sup> and if the decision accords with a four-part classification scheme.<sup>24</sup> An appendix sets out “Impact Assessment Levels” ranging from level 1, in which a decision will have “little or no impact on rights of individuals” including dignity and privacy, with impacts that are “reversible and brief,” to level 4, in which a decision

---

16 *CK v Magistrat der Stadt Wien* (CJEU, 27 February 2025), C-203/22.

17 *AI Act*, *supra* note 4, Arts 6(2), 14, 51, Annex III.

18 *Ibid*, Annex III, under heading 8, “Administration of justice and democratic processes,” and Recital 61.

19 See DADM, *supra* note 5.

20 See Guide, *supra* note 6.

21 On the legal status of DADM, see Teresa Scassa, “Administrative Law and the Governance of Automated Decision Making: A Critical Look at Canada’s Directive on Automated Decision Making,” (2021) 54:1 UBC L Rev 831 at 267–268, noting at 268 that “[d]irectives do not create actionable rights for individuals or organizations.”

22 DADM, *supra* note 5, cl 6.2.1

23 *Ibid*, cl 6.2.3.

24 *Ibid*, cl 6.2.1. See also Scassa, *supra* note 21 at 269–70, on the scope of the Directive’s application to persons external to government.

would likely have “very high impacts” on such rights that are “irreversible and perpetual.”<sup>25</sup> The appendix specifies that final decisions at levels 2 to 4 “must be made by a human.”<sup>26</sup>

The Directive also requires at all 4 levels that an affected person be provided a “general description of ... the criteria used to evaluate input [or client] data and the operations applied to process it,”<sup>27</sup> but it provides no further clarity on what this would entail. The Directive would, therefore, not rule out using a language model to render a decision accompanied by a general description of its underlying processes, but relegates AI to an assistive role.

In the fall of 2023, the Treasury Board released its “Guide on the Use of Generative Artificial Intelligence” for federal institutions.<sup>28</sup> Although it encourages the use of this technology in some areas, such as in determining whether a client is eligible for a service, it singles out ADM as a special case to be avoided.<sup>29</sup> It does so based on two concerns: 1) that the “design and function of generative models can limit institutions’ ability to ensure transparency, accountability and fairness” in decisions it makes or helps to make,<sup>30</sup> and 2) that using these tools to make “high-impact decisions” may violate a provider’s terms of use.<sup>31</sup>

Assuming that government employed its own language model to overcome contractual concerns, the Guide implies that the opacity involved in language model judgment is too significant to meet transparency and fairness concerns; that, in essence, the DADM’s requirements for a “general description” of underlying processes could either never be met where a language model decided a high-impact case or that any such description could never be sufficiently transparent to render it accountable and fair.

#### IV. New Capabilities of Generative AI for ADM

One inference to be drawn from law and policy in Europe and Canada on the use of ADM is that rights to a human decision and to an explanation are closely related to — though not limited to — concerns about the transparency and integrity of underlying algorithmic processes. Debate has continued as to whether the outcomes grounded in machine learning are sufficiently or effectively explainable, if not fully interpretable,<sup>32</sup> but the emphasis in the law and policy canvassed above is to some degree a function of the state of the technology at issue.

---

25 DADM, *supra* note 5, Appendix B

26 *Ibid*, Appendix C, citing cl 6.3.11.

27 *Ibid*, Appendix C, citing cl 6.3.2.

28 Guide, *supra* note 6.

29 *Ibid*.

30 *Ibid*.

31 The Guide cites as examples OpenAI’s prohibition on using their models for decisions about health, credit, employment conditions, law enforcement, and migration and asylum. It also cites Google’s prohibition on using their tools in “domains that affect material or individual rights or well-being.” See OpenAI, “OpenAI Usage Policies” (23 March 2023), online: <<https://openai.com/policies/usage-policies>>; Google, “Google Generative AI Prohibited Use Policy (2023), online: <<https://policies.google.com/terms/generative-ai/use-policy>>.

32 Zhao et al, “Explainability for Large Language Models: A Survey” (2023), online: *arXiv* <<https://doi.org/10.48550/arXiv.2309.01029>> [<https://perma.cc/NBA6-CDBS>]. See also Srihari Maruthi et al, “Language Model Interpretability — Explainable AI Methods” (2022) 2:2 AJMLRA; A Barredo Arrieta et al, “Explainable Artificial Intelligence (XAI): Concepts, Taxonomies, Opportunities and Challenges toward



Indeed, much of the existing law and policy is concerned with earlier forms of ADM involving algorithms or computational processes less complex than that found in generative AI and forms of AI that could not provide reasons for a decision. More recently, however, language models have begun to present a new possibility: automated decisions that contain extensive and responsive reasons to justify an outcome but rest on internal or underlying processes that remain largely opaque. This raises the question of whether the textual explanation a model provides, coupled with other safeguards such as testing for bias and general explainability, can satisfy concerns about transparency and impartiality in judgment, along with common law and constitutional requirements in relation to this.<sup>33</sup>

I defer this larger question to the next section and make the point here that scholars and jurists have begun to demonstrate these new capabilities of AI by conducting experiments with language models in the appellate, administrative law, and arbitration contexts.<sup>34</sup> The experiments demonstrate that large language models prompted to act upon party materials uploaded to the model (factums, briefs, written arguments) are capable of producing a reasoned opinion that is responsive to party-positions and engages in sophisticated legal analysis. Models can also make normative judgments and formulate new legal tests. In a critical assessment of these experiments, Grimmelmann, Sobel, and Stein concede that “LLM output now replicates rational, eloquent argumentation that applies precedent to novel facts. An LLM can produce text that may be formally indistinguishable from — or even formally superior to — the reasoning described by an opinion written by a human judge.”<sup>35</sup>

In other research, I have canvassed the methods and findings of three of these experiments, and describe two of them here briefly.<sup>36</sup> In 2024, Adam Unikowsky uploaded to Claude

---

Responsible AI” (26 December 2019), online: *arXiv* <<https://arxiv.org/abs/1910.10045>> [<https://perma.cc/82C3-E95N>].

33 On strategies to assess and mitigate bias in language models see the survey of scholarship in Yufei Guo et al, “Bias in Large Language Models: Origin, Evaluation, and Mitigation” (2024), online: *arXiv* <<https://doi.org/10.48550/arXiv.2411.10915>> [<https://perma.cc/L6VU-6AA8>]. See also Dong et al, “PRISM: A Methodology for Auditing Biases in Large Language Models” (2024) online: *arXiv* <<https://doi.org/10.48550/arXiv.2410.18906>> [<https://perma.cc/RHA2-RLHS>]. As to whether the opacity of underlying processes makes AI judgment different from human decision-making, see Eugene Volokh, “Chief Justice Robots” (2019) *Duke Law Journal* 68(6) 1135, noting at 1165 that “[i]f we are honest with ourselves, we often can’t really tell with confidence why we reached a particular judgment [...] We have reactions because of the real neural nets in our brains, and then we can offer explanations that we hope persuade”; see also Emily Berman, “A Government of Laws and Not of Machines” (2018) *Boston University Law Review* 98(4) 1277, noting at 1319 “contexts where human decision-making is itself opaque.”

34 Adam Unikowsky, “In AI we trust” (8 June 2024), online: *Adam’s Legal Newsletter* <<https://adamunikowsky.substack.com/p/in-ai-we-trust>> (accessed 7 April 2025) [Unikowsky, “Trust”] and “In AI we trust, part II” (16 June 2024), online: *Adam’s Legal Newsletter* <<https://adamunikowsky.substack.com/p/in-ai-we-trust-part-ii>> (accessed 7 April 2025) [Unikowsky, “Trust Part II”]. See also Kimo Gandall, Jack Kieffaber, & Kenny McLaren, “We Built Judge.AI And You Should Buy It” (January 28, 2025), online: *SSRN* <<https://ssrn.com/abstract=5115184>>; Eric A Posner & Shivam Saran, “Judge AI: Assessing Large Language Models in Judicial Decision-Making” (15 January 2025), online: <<https://ssrn.com/abstract=5098708>>; Robert Diab, “The Evolving Role of AI in Legal Judgement” (2026) *Law Innovation and Technology* 18(1) <<https://doi.org/10.1080/17579961.2026.2633688>>.

35 James Grimmelmann, Benjamin LW Sobel & David Stein, “Generative Misinterpretation” (18 June 2025, forthcoming in 63:1 *Harv J Legis*), online: *SSRN* <<https://ssrn.com/abstract=5309575>> at 51.

36 Diab, *supra* note 34.

3 court briefs from US Supreme Court cases in the current term. The model decided 27 of the 37 cases the same way the Court did and in the remaining 10, Unikowsky “frequently was more persuaded by Claude’s analysis than the Supreme Court’s.”<sup>37</sup> He documents Claude’s ability to formulate different and more elaborate legal tests in relation to those found in the actual decisions, and to apply these tests effectively to the facts in the case at bar.

In 2025, Gandall, Kieffaber, and McLaren conducted an experiment with software that two of them had created called Arbitrus.ai to demonstrate that it “performs just as well as human arbitrators right now” by giving it 100 hypothetical scenarios involving party materials created by another language model.<sup>38</sup> The authors found that “Arbitrus.ai performed as expected: it did not hallucinate; it answered the issues in controversy, and almost always — with two narrow exceptions — grounded those answers in relevant case law.”<sup>39</sup>

Once again, a language model can provide a rational basis for a decision, but not the technical or actual basis for why, in a given case, it decided in favour of one party over another. Models can thus offer a form of transparency, though the adequacy and value of this is contested.<sup>40</sup> Furthermore, the capabilities demonstrated in these experiments do not establish that language models are well suited to making decisions in any administrative or adjudicative context. However, further capabilities are on the horizon, including ADM involving auditory and visual input, giving rise to the possibility of hearings involving *viva voce* testimony assessed and decided by AI alone or in reliance on it to some degree.<sup>41</sup> These evolving capabilities collectively prompt a need to look beyond existing law and policy to address the question of whether common law and constitutional requirements of procedural fairness in Canada mark a hard limit on the possible use of these tools.

## V. Constitutional Limits on ADM

A legal decision made in full or partial reliance on AI would engage duties of procedural fairness at common law and rights under the *Charter*. In ways to be explored below, the *Charter* incorporates common law duties of procedural fairness, but it distinguishes between rights that apply when a person is “charged with an offence” (section 11) and rights that apply more broadly in situations involving “life, liberty, and security of the person” (section 7).

To address the first category, persons charged with an offence, rights in section 11 would seem to preclude both a verdict arrived at by fully automated processes or one made in significant reliance on them (i.e. helping to decide what findings of fact to make or how to apply

---

37 Unikowsky, “Trust Part II”, *supra* note 34.

38 Gandall, Kieffaber & McLaren, *supra* note 34 at 1.

39 *Ibid* at 56. The two exceptions involved the model drawing on cases with “disanalogous facts,” amounting to an erroneous reliance on certain earlier cases; errors, the authors point out, that a human judge might readily have made.

40 For critical views, see Grimmelman, Sobel & Stein, *supra* note 35; A Ebrahimi Afrouzi, “Robots, Thurgood Martian, and the Syntax Monster: A New Argument Against AI Judges” (2024) 37:2 Canadian Journal of Law & Jurisprudence 369; and John Tasioulas, “The Rule of Algorithm and the Rule of Law” (2023), online: SSRN <<https://ssrn.com/abstract=4319969>>. For a survey of contrary positions, see Robert Diab, “Vavilov and Generative AI” (2025) 61:3 Alta L Rev 1.

41 See the sources cited *supra* note 8.

law to facts). A right to a human decision can be inferred from the right in section 11(d) to a “fair and public hearing by an independent and impartial tribunal” and, in a non-military case where the jeopardy is five years or greater, from the right in section 11(f) to a “trial by jury.” An AI judge alone, or a judge acting in reliance on AI, may not offend the requirement that a hearing be “public” in nature by virtue of being capable of being conducted in an open and accessible manner, including to the media.<sup>42</sup> But a decision in either case would likely run afoul of the requirements for an “independent and impartial tribunal.” The law on judicial independence is extensive, but even a judge relying on AI in part to decide a case would fail to be independent, given the Supreme Court’s conception of this right. As Gontier J held in *Mackin v New Brunswick (Minister of Finance)*,<sup>43</sup> independence here refers to a “relationship between a court and others ... marked by a form of intellectual separation that allows the judge to render decisions based solely on the requirements of the law and justice.”<sup>44</sup> The Court has also defined impartiality in section 11(d) as “connot[ing] absence of bias, actual or perceived.”<sup>45</sup>

A verdict under section 11 resting partly on AI would raise concerns about a judge being impartial due to the possibility of their being affected by automation bias or the tendency to overestimate the value or reliability of the output of AI.<sup>46</sup> A verdict resting solely on AI would not overcome concerns about impartiality or independence due to the opacity of processes underlying an AI decision or recommendation, since these cannot be rendered fully transparent.<sup>47</sup> The lack of transparency would raise concerns that bias in a model’s training data may have affected the outcome, or that it was otherwise skewed by priorities or values in the design of the model’s functioning (a point canvassed further below).<sup>48</sup> And finally, the right to a “trial by jury” in section 11(f) would seem to require a trial decided by human peers, and this fact is not likely to change over the long term regardless of how much attitudes and opinions about AI in judgment might evolve.

Section 7 of the *Charter* provides an even broader basis for inferring rights to a human decision and to an explanation. Section 7 guarantees everyone “the right to life, liberty and security of the person and the right not to be deprived thereof except in accordance with the principles of fundamental justice.”<sup>49</sup> Where a decision engages life, liberty, or security of the person — e.g. in immigration or the provision of vital services — it must be made in a

---

42 See *R v Mentuck*, 2001 SCC 76 at paras 52–53 on the “public” component of 11(d).

43 [2002] 1 SCR 405.

44 *Ibid* at para 37.

45 *Valente v The Queen*, [1985] 2 SCR 673 at para 15.

46 Scassa, *supra* note 21 at 281, citing Linda J Skitka, Kathleen L Mosier & Mark Burdick, “Does Automation Bias Decision-Making?” (1999), 51 International Journal of Human-Computer Studies 991.

47 See the sources cited *supra* note 32.

48 Jonathan H Choi, “Large Language Models Are Unreliable Legal Interpreters” (28 February 2025), online: SSRN <[https://papers.ssrn.com/sol3/papers.cfm?abstract\\_id=5188865](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5188865)>. As to whether the opacity of underlying processes makes AI judgment different from human decision-making, see Eugene Volokh, “Chief Justice Robots” (2019) *Duke Law Journal* 68(6) 1135, noting at 1165 that “[i]f we are honest with ourselves, we often can’t really tell with confidence why we reached a particular judgment [...] We have reactions because of the real neural nets in our brains, and then we can offer explanations that we hope persuade”; see also Emily Berman, “A Government of Laws and Not of Machines” (2018) *Boston University Law Review* 98(4) 1277, noting at 1319 “contexts where human decision-making is itself opaque.”

49 *Charter*, *supra* note 9.



manner consistent with the “principles of fundamental justice.” The Supreme Court has held that fundamental justice in section 7 includes, at a minimum, certain common law principles of procedural fairness.<sup>50</sup> Specifically, “the tribunal which adjudicates upon [a person’s] rights must act fairly, in good faith, without bias and in a judicial temper, and must give him the opportunity to state his case.”<sup>51</sup> Beyond this, the requirements in a given administrative law case depend on the context at issue. The Court in *Baker* set out factors to determine what fairness requires in a given administrative context,<sup>52</sup> including the nature of the decision; the statutory scheme and direction set out in it; the importance of the decision to persons affected; “legitimate expectations” of persons challenging the decision; and any discretion open to the decision maker under statute about the process itself.<sup>53</sup> The *Baker* factors apply when determining what fundamental justice requires.<sup>54</sup>

As Hamish Stewart points out, while duties of fairness at common law are the same as those under section 7, a person would enjoy constitutional protection of these duties only where a decision engages their life, liberty, or security interests.<sup>55</sup> Where these interests are not engaged, a legislature may derogate from common law procedural requirements by express statutory language.

With this in mind, in what follows, I outline the contours in the administrative law context of a right to a human decision and to an explanation that would be rights at common law in cases that do not engage life, liberty, or security of the person, and constitutional rights under section 7 in cases that do engage one or more of these interests.

The literature on ADM and common law procedural fairness is extensive.<sup>56</sup> I seek to avoid redundancy with the scholarship to date by focusing here on how new capacities of language models to render reasoned decisions might be construed in relation to three rights in particular: the right to a hearing, the right to an impartial tribunal, and the right to reasons. I will argue that the limits imposed on the use of generative AI on each of these points in cases engaging section 7 interests outline the contours of a Canadian constitutional right to a human decision and to an explanation, and that these rights preclude full or even significant reliance on language models for ADM at present. In other cases, some of these requirements could be modified in legislation.

## The Right to a Hearing

This has two components: the right to know the case to meet and the right to be heard. On the first of these two components, the Supreme Court in *Charkaoui* held that “[t]he right to know

---

50 *Singh v Canada (Minister of Employment and Immigration)*, [1985] 1 SCR 177 at 212–13 [*Singh*]; *Pearlman v Manitoba Law Society Judicial Committee*, [1991] 2 SCR 869 at 883 [*Pearlman*].

51 *Singh*, *supra* note 50 at 213: Wilson J, in a concurring opinion, citing Fauteux CJ in *Duke v The Queen*, [1972] SCR 917 at 923. Justice Wilson’s holding in *Singh* is adopted by the majority in *R v Lyons*, [1987] 2 SCR 309 at 361.

52 *Baker v Canada (Minister of Citizenship and Immigration)*, [1999] 2 SCR 817 [*Baker*].

53 *Ibid* at paras 23–27.

54 *Suresh v Canada (Minister of Citizenship and Immigration)*, 2002 SCC 1 at para 115.

55 Hamish Stewart, *Fundamental Justice: Section 7 of the Canadian Charter of Rights and Freedoms*, 2nd ed (Toronto: Irwin Law, 2019) at 275.

56 See e.g. Gouge, *supra* note 1; Raso, *supra* note 1; Coglianese & Lehr, *supra* note 1; Scassa *supra* note 21.

the case to be met is not absolute,”<sup>57</sup> which means that a procedure might be fair in some cases where a judge hears from only one side.<sup>58</sup> In national security cases or where foreign confidences must be maintained, for example, privilege may preclude full disclosure of information a judge or decision-maker takes into consideration.<sup>59</sup> This is because of the way that broader “societal concerns” shape the “relevant context for determining the scope of the applicable principles of fundamental justice.”<sup>60</sup> Although regardless of context, where liberty or security interests are seriously affected, fundamental justice in section 7 cannot be satisfied without providing a “substantial substitute” for the privileged information.<sup>61</sup> This would entail “a summary of the information that is sufficient to enable the person to know the case to meet.”<sup>62</sup>

Scholars and policy-makers tend to agree that knowing the case to be met requires disclosure of the algorithmic or computational processes underlying an automated decision — sources of training data, variables considered, and the weight accorded to these variables, among other details.<sup>63</sup> The DADM, for example, requires notice of the fact that a decision will be made in whole or in part by an automated process and a “meaningful explanation to affected individuals of how and why the decision was made.”<sup>64</sup> One might provide a general summary of how AI decides, at the computational level, a case involving significant liberty or security interests (e.g. deportation). Yet there would be no compelling state interest in depriving an affected person of a more fulsome or accurate disclosure of these processes, as there was in a national security case such as *Charkaoui*.<sup>65</sup> Without more, the administrative expediency gained by resorting to an AI decision would not be a compelling enough reason to subject a person to such a process, given the stakes. However, where liberty interests are engaged to a lesser degree, such as in cases involving the issuance or extension of a visitor’s visa, fundamental justice might be satisfied by an AI process entailing limited transparency as to its inner workings if there were evidence that it decided cases consistently and without bias, and with other safeguards such as the availability of human review.

In relation to the other component of the right to a hearing, neither fundamental justice nor common law procedural fairness requires an oral hearing in every case. Section 7 will require an oral hearing where “a serious issue of credibility is involved” or where a person faces a significant period of detention,<sup>66</sup> though not in the extradition context.<sup>67</sup> Regardless of

---

57 *Charkaoui v Canada (Citizenship and Immigration)*, 2007 SCC 9 at para 57 [*Charkaoui*].

58 *Ibid*, citing *R v Rodgers*, 2006 SCC 15 and *Chiarelli v Canada (Minister of Employment and Immigration)*, [1992] 1 SCR 711.

59 *Charkaoui*, *supra* note 57 at para 58.

60 *Ibid*, paraphrasing *Ruby v Canada (Solicitor General)*, 2002 SCC 75, paras 38–44.

61 *Charkaoui*, *supra* note 57 at para 61.

62 *Ibid* at para 63.

63 See e.g. Marion Oswald, “Algorithm-Assisted Decision-Making in the Public Sector: Framing the Issues Using Administrative Law Rules Governing Discretionary Power” (2018) 376:2128 *Philosophical Transactions of the Royal Society* at 5, noting “the use of an opaque algorithm to generate an output informing a decision might obstruct the right to be heard if an individual is unable to understand the ‘gist’ of why the output was generated and so present an alternative case, or if the ‘learning’ nature of the tool makes it difficult or impossible to recreate the original decision.” See also Scassa, *supra* note 21 at 288.

64 DADM, *supra* note 5, cl 6.2.1.

65 *Charkaoui*, *supra* note 57.

66 *Singh*, *supra* note 50 at 213–14; *Charkaoui*, *supra* note 57 at para 28. See also Stewart, *supra* note 55 at 276.

67 *United States of America v Kwok*, 2001 SCC 18 at para 104.

whether a hearing is to be oral or in writing, it is unclear whether the duty to provide a hearing can be met without a human involved. Teresa Scassa queries whether there may be “something intrinsic to the concept of a hearing that requires a human to hear and to weigh the arguments made.”<sup>68</sup> As noted earlier, Canada’s DADM and both the EU’s GDPR and *AI Act* imply a positive response to this question in requiring a human to make a final decision or be “in the loop” in cases impacting individual rights to more than a minimal degree.

### **The Right to an Independent and Impartial Decision-Maker**

One possible reason to infer that a decision made partly or strictly in reliance on a language model would not satisfy a right to be heard is because this requires that the hearing be held “before an independent and impartial magistrate,” which includes a person who *appears* to be independent.<sup>69</sup> The irreducible opacity of a language model’s underlying processes, coupled with the impossibility of being sure an AI process is free of bias, would prevent an automated decision from at least appearing fully independent.

There are a host of reasons to doubt that a decision made in full or partial reliance on AI could be independent and impartial. A human judge relying to a significant degree on AI to make a decision would, as noted earlier, give rise to concerns about automation bias (and thus improper subdelegation) and being affected by possible bias in the training data. The Supreme Court’s decision in *Ewart v Canada* offers an example of concerns about both forms of bias in a risk assessment tool that resulted in a procedurally unfair process.<sup>70</sup> A language model designed and developed by a third-party might also reflect priorities and interests different from those of a court or tribunal using the model.<sup>71</sup> This would impugn the independence of any decision by allowing for the possibility of gaming an outcome by tailoring one’s submissions to those priorities and interests. A tribunal might overcome this by creating its own model, but this would not avoid the problem of bias in training data shaping the outcomes.

A further concern about both impartiality and independence of the language model as judge arises from evidence that what a model decides when asked to make a decision is affected by how a model is prompted. Jonathan Choi has demonstrated that “LLM judgments are highly sensitive to prompt phrasing, output processing methods, and model training choices,” leading to different outcomes and varying degrees of confidence in them based on slight changes in the wording of a prompt.<sup>72</sup> Yet Kieffaber and his co-authors have argued in response that Choi’s concerns apply to “off the shelf” language

---

68 Scassa, *supra* note 21 at 288.

69 Charkaoui, *supra* note 57 at para 29; see also Singh, *supra* note 50 at 212–13; Pearlman, *supra* note 50 at 882–83; *United States of America v Ferras*, 2006 SCC 33 at para 22. See also Stewart, *supra* note 55 at 276.

70 2018 SCC 30. The challenge in *Ewart* was not framed in terms of common law procedural fairness *per se*, but a breach of a statutory duty to ensure that actuarial tools used in prison risk assessments be accurate, and there was evidence in that case of a cultural bias in the tool’s application.

71 Maarten Buyl et al, “Large Language Models Reflect the Ideology of their Creators” (2024), online: *arXiv* <<https://doi.org/10.48550/arXiv.2410.18417>> [<https://perma.cc/88Q5-CESS>]; Jinhao Pan et al, “Beneath the Surface: How Large Language Models Reflect Hidden Bias” (2025), online: *arXiv* <<https://doi.org/10.48550/arXiv.2502.19749>> [<https://perma.cc/F39W-3AJY>].

72 Choi, *supra* note 48 at 1.

models, but not to bespoke models that incorporate forms of machine learning known as “classifiers.”<sup>73</sup>

In the legal context, classifiers can help models make more consistent judgments about factual scenarios by making decisions about them based not on next-token predictions alone but by identifying categories into which a factual scenario belongs. A debate continues to unfold as to whether language models can become consistent enough to be reliable, but the problem itself casts doubt on whether models can ever entirely overcome the possibility of being gamed — and if not, whether they could ever *appear* sufficiently independent and impartial.

### The Right to Reasons

The duty of procedural fairness at common law and fundamental justice under 7 require reasons to be provided in cases that engage significant liberty or security interests.<sup>74</sup> The conundrum to which language models give rise in their new capacity to decide cases based on party-materials is that they can provide a cogent and persuasive justification for an outcome, but they do not reveal the precise computational basis for reaching that outcome. This maps onto a distinction that Jennifer Cobbe makes between how a decision was arrived at (an explanation) and why it was made (reasons).<sup>75</sup> As Scassa puts it, “a right to an explanation explains how an algorithm operates to generate outcomes.”<sup>76</sup> Does either the requirement for reasons at common law or as part of fundamental justice under section 7 extend to a right to an explanation?

As noted earlier, Canada’s DADM requires persons significantly affected by automated decisions to be provided a “meaningful explanation” of “how and why the decision was made,” and comparable obligations can be found in EU law.<sup>77</sup> AI providers can discharge these duties by publishing information about a model’s operation,<sup>78</sup> but the question here is whether the textual reasons that a language model might provide would satisfy a common law or constitutional right to reasons — and in this respect we are in new terrain. The Treasury Board’s DADM and “Guide on the Use of Generative Artificial Intelligence” provide some indication as to what, under the test in *Baker*, it might be reasonable to expect about a process when deciding whether it is fair.<sup>79</sup> The DADM’s requirement for an explanation of both how and why an outcome was reached, along with the Guide’s resistance to the use of language models in judgment (given their opacity and possible bias), suggests that procedural fairness *should* include a duty to provide an explanation of underlying technical processes — in addition to any reasons a model might provide in its textual output. A model that only offers textual

---

73 Jack Kieffaber, Kimo Gandall, Steven Foster & Kenny McLaren, “LLMs are Bad Judges. So use Our Classifier Instead” (30 June 2025), online: SSRN <[https://papers.ssrn.com/sol3/papers.cfm?abstract\\_id=5331811](https://papers.ssrn.com/sol3/papers.cfm?abstract_id=5331811)>.

74 *Baker*, *supra* note 52; *R v REM*, 2008 SCC. As Stewart, *supra* note 55, notes at 287: “it is plausible to suppose that where interests as basic as those protected by section 7 are at stake, the principles of fundamental justice would require the decision maker to give reasons. Some cases strongly suggest as much, though they stop short of saying so explicitly.”

75 Jennifer Cobbe, “Administrative Law and the Machines of Government: Judicial Review of Automated Public-Sector Decision-Making” (2019) 39:4 LS 636 at 648.

76 Scassa, *supra* note 21 at 292.

77 DADM, *supra* note 5, cl 6.2.3; see also discussion of the GDPR and *AI Act* above.

78 See for example DADM, *supra* note 5, cl 6.2.3 and Article 50 of the *AI Act*, *supra* note 4.

79 *Baker*, *supra* note 52, at paras 23–27.

reasons why to accept a certain outcome as just would not entail a fair process because the model's reasons would not be the true or real reasons as to why it arrived at the outcome; the real reasons would be algorithmic, statistical, or computational.<sup>80</sup>

If the DADM and Guide can be read as a barometer of what a court would likely construe as necessary for a process to be fair, then the right to reasons at common law and under section 7 in certain cases would include a right to an explanation of processes underlying an AI decision. The view encapsulated in these policy instruments is not likely to change, at least in the short term. But this does not rule out the use of language models in administrative decision-making, including in high-impact cases. Over time, if language models become more consistent in operation and if underlying processes become more readily explainable, a model might be capable of meeting both the right to reasons and to an explanation. Though other concerns (impartiality, independence) might remain.

## VI. Conclusion

Much of the law and policy debate on the use of AI in legal decision-making predates the advent of generative AI. It took shape in response to more limited forms of AI and contemplated legitimate uses for them in ADM where their underlying processes could be rendered sufficiently transparent and impartial. New capabilities of language models to render reasoned decisions responsive to party submissions give rise to the possibility of automated decisions that are justified but not explained. Further new capabilities to hear and decide cases using auditory and visual techniques are on the horizon. This prompted a fresh look at common law and constitutional limits to using generative AI in ADM. Assumptions about the requirements of procedural fairness at common law and under the *Charter* help outline the contours of rights in Canada to a human decision and to an explanation that are, in certain case, constitutional rights, and rights likely to preclude relying on AI in whole or in part in high-impact cases for the foreseeable future.

---

80 Afrouzi, *supra* note 40; Tasioulas, *supra* note 40.



