

22 Can. J. L. & Tech. 201

Canadian Journal of Law and Technology

January, 2025

Article

Assistant or Authority? Artificial Intelligence in Business Decision-Making and the Director's Duty of Care

Margaret **Wilson**^{a1}

Copyright © 2025 by Thomson Reuters Canada Limited; Margaret Wilson

Abstract

The recent broad availability of Artificial Intelligence (“AI”) systems, particularly Large Language models, has become manifest in the corporate sphere in a range of applications. The startling capabilities of AI systems could establish them as members of corporate boards, and otherwise assist corporations in a number of ways. While the Canada Business Corporations Act (“CBCA”) precludes AIs from acting as board members, it does not restrict the use of AI by corporate directors. The duties of directors, however, and particularly the directorial duty of care as articulated under the CBCA, may affect the degree to which directors can utilize AI in business decision-making. This article discusses three AI-use scenarios to examine the degree of restraint that the directorial duty of care imposes on directors' use of AI in business decision-making.

Resumé

La récente large disponibilité des systèmes d'intelligence artificielle (<<IA>>), en particulier des modèles à langage large, est devenue manifeste dans la sphère des entreprises dans toute une gamme d'applications. Les capacités surprenantes des systèmes d'IA pourraient les établir comme membres des conseils d'administration des entreprises et aider les entreprises de plusieurs manières. Bien que la Loi canadienne sur les sociétés par actions (<<LCSA>>) interdise aux IA d'agir en tant que membres du conseil d'administration, elle ne restreint pas l'utilisation de l'IA par les administrateurs de sociétés. Cependant, les devoirs des administrateurs, et en particulier le devoir de diligence des administrateurs tel qu'articulé dans la LCSA, peuvent influencer sur le degré auquel les administrateurs peuvent utiliser l'IA dans la prise de décision commerciale. Cet article examine trois scénarios d'utilisation de l'IA pour examiner le degré de restriction que le devoir de diligence des administrateurs impose à l'utilisation de l'IA par les administrateurs dans la prise de décision commerciale.

***202 INTRODUCTION**

Alan Turing, the father of modern computer science, proposed what has become known as the ‘Turing test’ for artificial intelligence (AI) in 1950. A machine could be said to be intelligent, he suggested, when a human interlocutor could not distinguish written responses given by a human from those given by a machine.¹ Mustafa Suleyman, the cofounder of Google's AI branch, articulated an updated version of the test in 2024. He proposed that, in the modern context, a test for AI would assess whether the machine could act as an entrepreneur and run a business.²

While AI systems are not yet capable of the independent business actions that Suleyman suggested, the idea that powerful AI machines are capable of participating in business decisions is not purely hypothetical. In 2014, a Hong Kong company appointed an AI software, called ‘Vital,’ to its board of directors.³ The software was ostensibly intended to receive an equal vote and approval power on all financial business decisions. Vital was the first algorithmic program AI device to ‘sit’ on a board but is no longer unique. More recently, a Chinese video-game corporation appointed an AI to serve as an executive officer.⁴

Outside of appointments to the C-suite, the use of AI systems as authoritative sources in business decision making has also increased. For example, Salesforce executive, Marc Benioff, has used an AI device called "Einstein" during meetings and in performance evaluations.⁵ Some commentators have predicted that AI devices like Einstein will soon "become core to decision-making in the boardroom."⁶

The appointment of an AI -- an amorphous concept in itself -- to a directorial position, as in the Vital example, is not permissible under Canadian ***203** corporate law. Under the federal *Canada Business Corporations Act (CBCA)*, the focus of this article, a director in Canada must be "an individual," precluding the appointment of non-human directors.⁷ Similar provisions also exist in many provincial business corporation statutes.⁸

The *CBCA* does not specifically restrict the use of AI by directors. However, the duties of directors, as articulated under the *CBCA*, could affect the degree to which directors may utilize AI in their business decision-making. The use of AI by directors raises provocative questions around the directors' duty of care and the delegation of decision-making tasks within that duty. This article will discuss three AI-use scenarios to examine the degree of restraint that the directorial duty of care imposes on directors' use of AI in decision-making in *CBCA* corporations.

Part I will summarize the director's duty of care under the *CBCA*. Part II will engage in a technical and theoretical discussion of AI, using Large Language Models (LLMs), like ChatGPT, as an illustrative example. Part III will analyse three AI-use scenarios and discuss the implications of each for the directorial duty of care in the use of AI in business decision-making.

I. LEGAL BACKGROUND

(a) The Director's Duty of Care

Directors oversee the business of a corporation. They serve on the board of the corporation and make decisions about the direction and future of the corporation. They are involved in the appointment of officers and other key managers of the corporation's affairs.⁹ In carrying out these roles, directors are subject to several connected duties, including the duty of care. To understand how AI interacts with the directors' duty of care, it is important to have a clear understanding of that duty as articulated under the *CBCA*.

The directorial duty of care developed at common law. In *Re City Equitable Fire Insurance Company Ltd*, Romer J. articulated the duty of care as follows:

A director ... does not owe a duty to his company, to take all possible care. It is some degree of care less than that. The care that he is bound to take has been described ... as "reasonable care" to be measured by the care an ordinary man might be expected to take in the circumstances on his own behalf.¹⁰

***204** This articulation of the duty reveals a standard of care that is partly subjective and partly objective. The codification of the duty of care and subsequent statutory reforms "reflect its fundamental wisdom."¹¹ The common law duty and standard of care were codified in the *CBCA* as follows:

Every director and officer of a corporation in exercising their powers and discharging their duties shall ... exercise the care, diligence, and skill that a reasonably prudent person would exercise in comparable circumstances.¹²

The statutory duty requires more of directors than the common law test. Per the Supreme Court of Canada (SCC) in *Peoples*, the statutory duty has a largely objective standard of care.¹³ Components of subjectivity remain, for example, as regards the 'skill' of the director. The director is required to be careful in carrying out tasks, but not more so than his or her skills allow. However, an unskilled reasonable person might be expected to prudently recognize their limitations and seek advice from a more skilled person.¹⁴ A director may be found to have breached the standard of care if they failed to act in a manner as described above.¹⁵ In practice, the statutory duty also requires directors to devote time and attention to their directorial role.¹⁶ However, the content of the duty -- the care requirement -- remains largely unchanged between the common law and the statutory duty. It "primarily requires some degree of attention to detail in performing any particular managerial task."¹⁷ Note also that directors cannot contract out of the duty of care. Though they may delegate some tasks, directors must retain the final authority in decision-making.¹⁸

The duty of care is owed primarily to the corporation. However, unlike the duty of loyalty, directors may also owe a duty of care to creditors of the corporation and to other stakeholders.¹⁹ Directors must consider how the *205 interests and expectations of stakeholders might be affected when they make business decisions on behalf of the corporation.²⁰

(b) The Business Judgement Rule and Due Diligence

When the duty of care is breached, the Business Judgement Rule (BJR) serves to provide judicial deference to directorial decisions. Like the duty of care, the BJR developed at common law as a reflection of the discomfort of the English common law judiciary in passing judgment on the appropriateness of directors' exercise of business judgment. Business judgments and decision-making are outside of the expertise of the judiciary but are precisely the role of directors.²¹ The BJR functions to provide judicial deference to the expertise of directors in the conducting of business affairs. The rule is intended to incentivise some risk-taking by directors, and the smooth functioning of the market.²² It shields directors from liability when a business decision was not perfect but was made in good faith.²³ The BJR has been adopted in Canadian jurisprudence and was articulated by the SCC in *Peoples*:

In determining whether directors have acted in a manner that breached the duty of care, it is worth repeating that perfection is not demanded. Courts are ill-suited and should be reluctant to second-guess the application of business expertise to the considerations that are involved in corporate decision making, but they are capable, on the facts of any case, of determining whether an appropriate degree of prudence and diligence was brought to bear in reaching what is claimed to be a reasonable business decision at the time it was made.²⁴

To benefit from the BJR, directors must demonstrate that their decision was "fully informed."²⁵ This may include demonstration by directors of a clear *206 understanding of their duties, receipt and review of information relevant to the decision at issue, time spent on the review of relevant documents, and compliance with appropriate corporate procedures.²⁶

Directors may also be shielded from liability resulting from a breach of the duty of care by demonstrating due diligence. Ss.123(4) and 123(5) of the *CBCA* provides relief from liability where an outside expert has been relied upon.²⁷ Due diligence in the exercise of care may include the good faith reliance on corporation financial statements or a report from an outside expert, “whose profession lends credibility” to their statement.²⁸ This is usually taken to mean accredited professions, like lawyers and accountants.²⁹

II. TECHNICAL BACKGROUND

(a) Operation of Artificial Intelligence

In order to analyse the interaction of the directorial duty of care with AI, a grounding in the technical operation of AI systems is essential. AI is a specialized discipline of computer science.³⁰ It encompasses a range of applications and techniques, including fields like robotics, computer vision, and natural language processing. As such, a definition of AI, per se, is challenging. Exact definitions change with the field and application. For the purposes of this article, and in keeping with Turing's 1950 test, AI may be defined as “a diverse set of computational procedures or techniques (such as machine learning algorithms) that, based on data, perform tasks to varying degrees of autonomy that would be considered intelligent if performed by humans.”³¹ The goal of the discipline is to “simulate human intelligence in a way that enables the replacement and/or augmentation of human cognitive tasks by computer processes.”³²

***207** The foundational approach to current AI systems is the use of machine learning (ML) techniques. ML algorithms are the engines that drive many AI systems, including popular systems like ChatGPT. ML is important to AI systems because an ML system can develop capacities beyond those which are encoded. The capabilities of the machine are not solely dictated by the written code program. An ML system is, rather, a predictive tool. Using data, sourced from a broad variety of datasets, the program can develop and build a hypothesis. Using the hypothesis as a model, the program can make “predictions about the value of some unknown variable based on one or more known variables”³³ and use the prediction to solve the problem.³⁴

All ML algorithms operate in relation to a “predictive goal.”³⁵ This is an abstract goal that has been “translated into a decision about, conceptually, what to predict.”³⁶ The goal is expressed by the metric the model will use to achieve this predictive output and is closely tied to the intended use of the system. Often the goal is the accuracy of the prediction or estimate of some variable.³⁷ The variable of interest changes with the application of the algorithm.

The availability of training data is also key to the operation of these systems. Training data, “in essence a database of facts about known cases,”³⁸ is used to formulate the prediction that the machine makes. These datasets are often substantial, since a good prediction will require a large snapshot of outputs that the program is intended to predict. Based on the training set, a hypothesis can be developed from which the prediction follows. Training data places limitations on the predictions that result from ML models because predictions are based on data. Predictions are only as good as training data: ‘garbage in, garbage out,’ as the saying goes. Data carries with it implicit biases, cultural assumptions, and more. The provenance of data can be a source of bias, sometimes in unexpected ways. For example, some AI systems learned to write ‘work’ correspondence by being trained on emails obtained and used as evidence during the Enron bankruptcy investigation.³⁹ The systems were learning exclusively from the ***208** communications of “people who thrived at a morally compromised US corporation.”⁴⁰ This creates concern about the biases implicit in these communications, which the machine has inherited.

ML models are based on three primary statistical approaches. The first, regression modelling, develops a statistical trend on the basis of the training set. New, or unseen, data is then placed and fitted to the trendline to make a prediction.⁴¹ The exact type of regression model used (such as linear, parametric, non-parametric, Bayesian models, etc.) varies depending on the size and complexity of the training dataset. The second statistical approach is a decision-tree method. In this approach, a starting event is “filtered through a complex ‘if-then’ analysis”⁴² in sequence to make a prediction. Both the regression and decision-tree

procedures are limited to relatively simple predictions which are entirely mathematical in nature. They cannot make connections between, for example, two decision-trees or two regression models, unless explicitly coded to do so. Further, these models are limited by the rigidity of their structure. All input data is reduced to a binary 'yes-no' that then informs the outcome. While they can extrapolate meaningful predictions from unseen data, application of these models is limited to relatively simply logic patterns.

The third ML approach, the neural network, seeks to address these limitations. The neural network methodology is patterned on the computational approach taken by the human brain and is used in a range of AI tasks and systems. In the neural network methodology, individual 'nodes' perform specific computational tasks. These are usually relatively simply mathematical operations, like matrix multiplication.⁴³ The network links the individual nodes together, via "synapse-like links that have adjustable weights."⁴⁴ The outcome of each computation feeds forward to the next node, and informs the overall prediction made by the system in the final output. The systems can also contain loops, such that the outputs from various units are fed backwards as input data, and information flows through the system. The inclusion of loops allows "the signal values within the network [to] form a dynamical system that has an internal state or memory."⁴⁵ This linkage and network effect allows for a more nuanced approach to predictions because *209 connections can be made between various factors. Further, the combination of linkage networks leads to the development of 'deep networks,' where layers of neural networks contribute to increasing computing power.⁴⁶ In deep networks, each layer of the network can help other layers to learn, increasing the accuracy of the prediction.

Within these approaches, there are also different forms of learning algorithms that enable the machine to learn from data and make accurate predictions. The most regularly used algorithmic approaches are supervised learning, unsupervised learning, and reinforcement learning.⁴⁷

Supervised learning algorithms are taught by being initially supplied with "correct" answers. Training sets have labelled outcomes which represent the correct values predicted from the associated input.⁴⁸ For example, suppose a collection of images is supplied to the algorithm. It learns, through labels on the image, which pieces of the collection contain a cat. Having processed this training data, the algorithm is then able to make accurate predictions on the appropriate label for unseen data.⁴⁹ As per the example above, the system is able to identify whether the image should be labelled as including a cat. In essence, systems using this learning algorithm can extrapolate a generalized prediction from the training data.

Unsupervised learning algorithms do not start from a point of understanding the 'correct' answer. Rather, they are intended to find structure in unlabelled data. These algorithms identify common features within the training data set and groups these features together in clusters of data. When dealing with new data, the algorithm assigns that data in relation to established clusters of data in the existing data set. Good examples of supervised learning algorithms are those which are used to produce 'recommended' content on streaming and social media platforms. The recommendations are based on content previously "liked" or consumed by the user.

The third type of ML algorithm is a reinforcement learning algorithm. In this approach, the system learns through feedback, -- reinforcement -- trial and error, and progress towards the system goal, "which of the actions prior to the reinforcement were most responsible for it, and to alter its actions"⁵⁰ to achieve an outcome closer to the overall goal. The feedback of the system is evaluated to *210 inform the next output. Examples of such systems include those which are adept at playing games, like chess. With each move, the system assesses whether it is closer to winning.

It is from these learning algorithms, particularly in unsupervised and reinforcement learning, that AI systems develop their 'black box' quality. Because the algorithm is learning with each piece of data or each use, how the machine is arriving at the predictions is not clear to the user, operator, or even the programmer. With such a variety of factors informing a prediction, a user will have very little understanding of how, for example, "a deep network computes its output from its inputs."⁵¹

To tie these concepts together, it is useful to consider an existing, popularly used LLM AI system: ChatGPT. The system comprises three key components, which can be readily scaled up to process and use large quantities of data. First, the system is trained on language models through supervised learning. The system is provided with as much data as possible, with a goal of accurately predicting the next word in a given sentence. The more exposure to language, the more effective the system becomes at producing sensible, coherent predictions. In the second stage, supervised learning is again used to fine-tune the program to produce answers to questions. This stage is similar to the first, except that the output is more than a single word. The third stage enables the system to respond to a user prompt. This component of the program uses reinforcement learning from human feedback, allowing the user to indicate preferences on the response the system output. The feedback helps to improve the system results.⁵²

One of the issues that has been noticeable in ChatGPT is the production of hallucinations: the generation of nonsensical or inaccurate “responses unsupported by their training data.”⁵³ The hallucinations that ChatGPT produces derive from the first and second stages of the process. The overall goal of the system is to accurately predict the next word, or sentence. The goal is only to give answers, so the systems “prioritize generating the best possible guess in their responses” rather than the production of correct answers.⁵⁴

***211 (b) Considering Artificial Intelligence**

It is tempting to use ‘action’ words (like ‘thinking’ or ‘writing’) when referring AI systems in operation, because they do appear to be completely autonomous. The outcomes of some generative AI systems, for example, are impressive and can make it appear that there is a genuine, creative, synthetic intelligence behind the output. But, as the foregoing demonstrates, AI systems are fundamentally statistical. They are not ‘thinking’ in the way that a human or other non-artificial intelligence does. Rather, these systems “run a simulacrum of thought,” without consciousness or awareness behind the thinking.⁵⁵

The extent to which AIs approximate thought becomes more apparent when considering what kind of logic patterns the systems are performing when they produce answers. To illustrate the limitations of currently existing AI systems, per Larson, it is useful to examine three key modalities of thought and reasoning: deduction, induction, and abduction.⁵⁶

In deductive reasoning, a fact is inferred when other known premises are present. AI and computer systems are good at deductive reasoning, because it “supplies a template for “perfect” and precise thinking for humans and machines.”⁵⁷ This logic pattern is relatively easy to program and simulate using data. If a given premise is present, a true conclusion inevitably follows.⁵⁸

Inductive reasoning is the acquisition of knowledge from experience, which forms the basis for conclusions.⁵⁹ The statistical nature of AI means that the systems also utilize inductive reasoning, because it is “tied inextricably to data and frequencies of phenomena in data.”⁶⁰ The data that the AI system is trained on -- and, by extension, the knowledge that it has acquired from ‘learning’ -- is extended to draw conclusions about ‘unseen’ data. The inductive reasoning of AI systems is, however, limited by data it has already been exposed to, since “we can't infer what we don't know.”⁶¹

***212** The final mode of thinking is abduction. Abductive reasoning derives an explanation for a fact from a hypothesis that is linked to the circumstances. This type of thinking is “tied closely to reasoning from events to their causes.”⁶² Essentially, abductive reasoning is a kind of educated guessing, the generation of a reasonable hypothesis. A human confronted with a situation will produce a reasonable or plausible guess as to what happened to cause the situation before them, based on their knowledge of circumstances and common sense. Abduction is used in creating algorithms but is not a capability that the machine can replicate itself.⁶³ AI systems are relatively successful with deduction and induction because these types of reasoning suggest that something ‘must be’ or ‘is likely to be’ true, given certain data and premises. Abductive reasoning by contrast suggests that something “may be” true.⁶⁴ This lack of certainty is confounding for AI reasoning systems. Further problematic for AI systems is that the nature of facts -- or data for the machine -- change when using abductive reasoning. As Larson notes,

“[w]hereas induction treats observation as facts (data) that can be analyzed, abduction views an observed fact as a sign that points to a feature of the world.”⁶⁵

Abductive reasoning is challenging for AI system, in part because such a reasoning pattern is an inherent selection problem of choosing the most reasonable explanation from all possible explanations. To successfully abduce something “we must solve the selection problem among competing causes or factors, and to solve this problem, we must somehow grasp what is relevant in some situation or other.”⁶⁶ An AI is not capable of such a reasoning approach. It literally cannot ‘think outside the box.’

Abduction is also the kind of reasoning that allows for the development of new ideas. An AI system is precluded from producing anything of genuine novelty by its training and dependency on data. The fact that the machine is ***213** using data for decision-making means that “the knowledge available to the system is, in some very real sense, already discovered and written down.”⁶⁷

(c) The Reasonable Artificial Intelligence?

To abduce, then, is to have an ability to be selective, contextual, and to have a true understanding of circumstances. The ability to abduce reveals an ability to be reasonable. Reasonability, in law, requires some aspect of subjectivity. The “reasonable person” will vary with the circumstances at issue, requiring accounting for contextual and subjective factors.

Even using the logic patterns that AI systems are relatively proficient at, the systems are not able to be sensitive to context. The inductive abilities of AI systems are limited in this way because they lack ‘common sense’ knowledge. Common sense requires a “rich understanding of the real world.”⁶⁸ The AI system only ‘knows’ the data set that it has been shown and trained to optimize, but with no genuine, contextual understanding of that data that comes from knowledge of everyday concepts.⁶⁹ Conceivably, an AI might be able to acquire some commonsense understandings through learning from data. But since commonsense refers to a broad base of knowledge that varies with place, culture, and time, it would be very challenging -- and take a lot of data -- for an AI system to acquire such an understanding. Additionally, adding more data to the system is not necessarily a viable solution. It has not, for example, proven to be a solution to the limits of inductive reasoning in AI. In fact, with the addition of more data, “systems ha[ve] more and more chances to get everything wrong.”⁷⁰

The lack of context with which AI processes information can also have a moral quality. For example, the systems take data at face value without a sense of history or background. As Smith observes, we do not know “what normative standards, registrations schemes, ethical stances, epistemological biases, social practices, and political interests have wrought their influence across the tapestry” of the data.⁷¹ Further, current AI systems are not designed to assess the ‘truth’ of a statement. As discussed above, the goal is to give a statistically likely answer to ***214** the user. The fact that the AI system has an operational goal precludes understanding of what it is processing. Turing’s original test alludes to this concept. To successfully pass the Turing test, an AI system doesn’t need to comprehend what it is reporting. It merely needs to give a sufficiently convincing impression of understanding. The information that this requires is limited: “a knowledge base to draw from and make conversational connections and an appropriate grasp of syntax.”⁷² Because of this, AI systems have been analogised as “‘stochastic parrots,” since they repackage and repeat information, with no true understanding of what they are saying.⁷³ Smith illustrates this idea succinctly: “When our iPhone tells us that there has been an accident on the highway, [or] a medical system claims that nausea affects 32.7% of patients who have taken a given drug ... The fact that these systems use natural language can obscure to us the fact that they understand neither traffic nor nausea ...”⁷⁴

The machines do not care about or understand the outcome, except inasmuch as it meets the operating goal. Otherwise, “they don’t give a damn”⁷⁵ about the actual outcome itself, nor the consequence of the outcome produced.⁷⁶ Such detachment is distinct from impartiality. There is much discussion of how AI, and computers generally, make ‘better,’ more “neutral’ decisions in comparison to humans, because the system has no horse in the proverbial race, and is therefore impartial. But impartiality requires some level of awareness of the factions at play and judging where the balance lies.

AI systems also lack subjectivity. Subjectivity requires some level of understanding of a scenario. Russell and Norvig illustrate this with a simple scenario. They note that machines have no ability to be subjective, since they cannot “use their own subjective apparatus to appreciate the subjective experience of others ... if you want to know what it's like when someone hits their thumb with a hammer you can hit your thumb with a hammer. Machines have no such capability.”⁷⁷ Subjectivity involves putting oneself in another's position, though usually less physically than in the example above. An AI system is not capable of this kind of judgment. It is not capable of judging itself or why it ***215** acted in a certain way to produce an outcome, except in statistical terms and through inferences developed through ML.⁷⁸ This becomes even more apparent, when questions requiring more nuanced answers are put the system. For example, AI systems have difficulty distinguishing between good and bad actions. One researcher demonstrated that when asked to identify ‘good’ or ‘bad’ actions in fairy tales, a ML system was able to place characters into clusters of similar actions but was not able to conclusively label the characters or their actions as ‘good’ or ‘bad’. Given an understanding of the statistical nature through which AI systems produce answers, this is hardly surprising. Answering questions about whether something is ‘good’ or ‘bad’ requires subjective knowledge and an understanding of the contextual locus of values.⁷⁹

Because of the lack of ability to reason contextually, subjectively, or to make reasoned guesses, an AI system is not capable of being ‘reasonable’ in the legal sense of the word. Rather, in the current developmental stage, AI systems possess exclusively ‘reckoning’ abilities: “types of calculative prowess at which computer and AI systems already excel -- skills of extraordinary utility and importance.”⁸⁰ What they lack, therefore, is judgment: “dispassionate deliberative thought, grounded in ethical commitment and responsible action, appropriate to the situation in which it is deployed ... a form of thinking that is reliable, just, and committed -- to truth, to the world as it is.”⁸¹ This is precisely the kind of thought that is expected of someone who is supervising, managing, or exercising care.

Concerns arise in the use of AI because the system appears to be operating “at the upper end of the knowledge pyramid, but, in reality, it is only performing without its own understanding what the human mind has included in the context that the artificial intelligence has been trained with.”⁸² This reflects the socially pervasive idea, alluded to earlier in this section, that AI systems, and computers generally, make ‘better’ or ‘neutral’ decisions compared to human decision makers. The concern is, then, that ‘the human in the loop’ -- to borrow the phrase beloved by computer scientists -- “will rely on reckoning systems in situations which require genuine judgment.”⁸³

***216** This concern will be illustrated in the context of business decision-making in the scenarios that follow.

III. THE DIRECTOR AND ARTIFICIAL INTELLIGENCE

(a) The State of Artificial Intelligence Use in the Corporation

The use of AI systems by corporations is accelerating, with a recent survey suggesting that two-thirds of global businesses view AI systems as critical to business success.⁸⁴ Currently, AI is used by businesses in a range of decision-making areas including “due diligence of mergers and acquisitions, profiling investors, auditing annual reports, validating new business opportunities, analyzing and optimizing procurement, sales, marketing, and other corporate matters. Most businesses are already utilizing some form of AI, algorithms, and various platforms, such as ChatGPT.”⁸⁵

As discussed above, AIs are most effective where there are large amounts of data to learn from and to process. In the corporate context, AI systems are therefore likely to be used “where the most complex decisions need to be taken within corporations”⁸⁶ because such situations require the consideration of large amounts of data and many interrelated factors. Moslein suggests that “strategic business decisions ... are concurrently the most complex ones that have to be taken within corporations.”⁸⁷ Because AI systems could be used in decision-making, the directorial duty of care is implicated. It is thus important that the protocols regarding AI use be thoroughly delineated.

To do so, three scenarios will be used to examine the interplay between the statutory duty of care and the use of AI by directors in decision-making. Director A is the protagonist in each scenario. Director A is a director of a large corporation: XYZ Corporation. He does not have a background in a recognized expert profession and is not expected to have skills beyond those of an experienced businessperson. The decision that Director A must make in each scenario is a strategic business decision that will have significant ramifications for the future of the company. As such, he has been provided with relevant internal *217 financial statements and reports from qualified external experts. In each scenario, the AI technology that Director A has elected to use is ChatGPT.

Note that, although this analysis will proceed by considering the duty of care in isolation, use of ChatGPT would likely have broader ramifications for the director. For example, it might constitute a breach of other connected directorial duties under the *CBCA*, like confidentiality and loyalty. Because ChatGPT is a publicly available AI tool, once a director places information about the decision into the tool, it is 'known' to the system and, thus, no longer confidential.

Note, too, that these scenarios are necessarily limited by their simplicity. For example, in most large corporations, an individual director does not have much decision-making power. Decisions would be made by vote in consultation with the board of directors. There is no requirement for directors to justify their vote to one another.

(b) Scenario 1: Documenting the Decision

In this scenario, after review of expert reports and internal financial statements, Director A has come to a decision. He prompts ChatGPT to write a discussion and justification of his business decision.

Director A has met his duty of care in this instance because he has demonstrated reasonable prudence in arriving at the decision. By reading the financial statements and expert reports, Director A has become informed about the circumstances in which he must make the decision, the options available to him, and the projected outcomes of those decisions.

Part of the reason that this decision does not constitute a failure to meet the duty of care is because Director A does not significantly delegate his reasoning in arriving at the decision. The only point of delegation is in relying in good faith on the expert reports, which is permitted by the *CBCA* under ss.123(4) and (5).⁸⁸ There is no substantial delegation of reasoning to the AI system. Director A makes the decision autonomously and uses ChatGPT as a supporting tool. In this scenario, the AI system is in the role of an assistant, since the "decision rights remain solely with the human being."⁸⁹ Writing a justificatory discussion is largely an administrative task, with very little judgment involved. However, Director A should review the written discussion of his decision produced by ChatGPT to ensure that it is sensible and accurately represents his thought processes in arriving at the decision.

If, in the future, the decision was to have negative ramifications for the corporation or other stakeholders, the BJR would likely apply to insulate the decision from substantive review by the courts. The decision likely falls under the BJR, because it is within the range of reasonable options for business decisions.

***218 (c) Scenario 2: Supporting the Decision**

In this scenario, Director A has reviewed the expert reports and internal financial statements provided to him. He finds himself 'stuck,' and unsure of what decision to make. To help, he prompts ChatGPT to suggest a decision. He uses the chatbot's output to aid his thinking and, with further consultation of relevant documentation, arrives at a decision.

Up to the point of consulting ChatGPT it is likely that Director A has met his duty of care in this scenario. His review of the documentation is appropriate. Further, the fact that he is 'stuck' indicates a thoughtful and prudent approach to the decision and appropriate course of action. He is not being hasty, but rather is exercising thought and care in the decision.

The consultation of the AI system marks the point at which the analysis becomes less clear. As long as the AI is used to prompt additional thinking on the part of Director A, this use will not constitute a breach of the duty of care. Director A should robustly

scrutinize the decision proposed by AI and use it as a stepping off point for additional thinking. He should consider whether he concurs or disagrees with the decision that the AI produces and the reasons for his thinking. Director A should also be alert to the fact that ChatGPT is not concerned with truth or ethical decision-making. As discussed above, it is only intended to calculate statistically likely written responses.

There is more concern with the use of AI in this scenario because this use of ChatGPT represents a more substantial delegation of reasoning. This use of AI in decision-making goes beyond that of purely assisting. The AI serves to augment and enhance directorial thinking. Commentators have suggested that this paradigm of AI use, where “humans and machines share decision rights and learn from each other,”⁹⁰ will be the most common use of AI in the corporate boardroom context. It is imagined that in this situation, AI systems “will complement their human counterparts, allowing decision-makers to free up their minds and their time to think about larger strategy.”⁹¹ This points to the observations made about Director A's conduct above. He must still retain the essential decision-making authority. The AI must still occupy the role of assistant. This is within its capabilities. There is no expectation of reasonableness for the AI in an assistant role. Delegation of responsibility beyond this would result in substantial decision-making authority being conferred on the AI. As Moslein notes, “human directors may delegate decision rights to machines, but they must always keep the ultimate management function for themselves.”⁹² ***219** Failure to do so is likely to constitute a breach of the duty of care. Complete delegation of reasoning in decision-making is not appropriate in part because AI systems have no capacity for reasonableness. The director would not be acting in a reasonably prudent manner by delegating to the AI and would not be fully informed about the decision.

If the decision that Director A makes were to be problematic in future, the BJR would likely apply, as it did in Scenario 1, to shield Director A from liability. Although this decision was AI-assisted, Director A was still ultimately responsible for it. He made the decision in an informed manner from a set of reasonable alternative decisions.

(d) Scenario 3: Making the Decision

In this scenario, Director A has not reviewed expert reports of the internal financial statements in detail. With some preliminary information, he prompts ChatGPT to devise a business decision for him. He reads the chatbot response, finds it to be acceptable, and adopts the output as his decision.

Director A is not likely to have met his duty of care in this scenario. Director A has not been diligent, nor has he acted with reasonable prudence in arriving at this decision. There has been no consideration of alternative decisions, nor has the decision been made in a manner which is fully informed. Director A, in fact, has very little information since he has not substantially reviewed the expert documentation provided. Further, Director A has not substantially examined the decision proposed by ChatGPT. Given the widespread coverage of ChatGPT's tendency to produce hallucinations, a reasonable person would likely review the AI-proposed decision carefully before adopting it.

The duty of care is not met in this scenario in part because of the substantial delegation of reasoning to ChatGPT. This scenario is an example of a director conceding much of their decision-making rights to the machine. While the delegation of reasoning to the AI is not total, the reasoning that the AI is simulating represents much of the thought going into the decision. The AI is no longer an assistant. It has been recognized as an authority.

This is a problematic position for an AI to occupy because these systems are not capable of the kind of reasoning that the law expects directors to undertake in making a decision. An AI, alone, is not capable of being reasonable within the meaning of the duty of care. As discussed above, the systems cannot abduce and thus cannot make appropriate selections between alternatives. An AI cannot make an ‘educated guess’. A ‘reasonable’ decision is a true operation of judgment and abduction. The work involved in making a business decision is “judgement work.”⁹³ Making a business decision requires, among other things, “problem ***220** solving and collaboration, strategy and innovation, and relationship with individuals and stakeholders” and other factors, like commonsense.⁹⁴ As discussed in the previous section, an AI is not capable of true ‘judgment’ calls or commonsense thinking, since such choices require an understanding of, and investment in, the circumstances of the decision. Because of this, the substantive delegation of decision-making authority to AI systems is not appropriate.

The AI as director embodies the idea that the AI system is making a decision that is 'neutral' when compared with the decision made by a human director. Directors themselves are not 'neutral' decision makers; they make decisions in the context of personal views and biases. The allure of the AI in decision-making is that it, theoretically, comes without human prejudices. As Hildebrandt notes, "[t]he aura of objectivity that is often attributed to computing systems may have a strong influence on the human decision maker."⁹⁵ But, as discussed above, the neutrality of AI systems is illusory. AI systems make decisions with biases that are inherent in their training data. Directors, while not 'neutral', do not make decisions in an analytical vacuum like AI systems do. Directors have a better understanding of the range of possible acceptable decisions than an AI. The director has an understanding -- and a duty -- of reasonability.

Moslein suggests that, because of this tendency to view AI systems as neutral, there are two instances in which an AI-as-authority is likely to emerge. He suggests that "machines ultimately take over all decision rights, either because humans increasingly trust the machines' abilities to decide, or because decisions have been taken so quickly or require so much data that humans are simply unable to decide."⁹⁶ In such situations, directors should be aware that systems are not without fault, and AI system outputs should be monitored accordingly. Such heightened diligence will help to ensure that the duty of care is met. Directors should be conscious that a decision is not a 'better' one because it comes from an AI system. Simply because an AI system is being used, does not mean that that director can walk away from the decision-making process: "[c]orporate management retains the duty to ensure that automated tasks are performed properly."⁹⁷

In this scenario, the BJR will likely not apply to protect the director from liability arising from a breach of the duty of care. While a court might not scrutinize the substance of the decision itself, a judge could certainly conclude that, in the circumstances, by relying unquestioningly on the decision of an AI, *221 Director A did not bring "an appropriate degree of prudence and diligence"⁹⁸ to the decision.

Some commentators suggest that the BJR will be essentially obsolete if AI decision making were present. This proposition is made on the basis that because the AI has no vested interest in the decision, the BJR no longer serves its original purpose of incentivising risk and protecting risk-takers, within reason.⁹⁹ However, there are good policy reasons for retaining the BJR in the age of AI in the boardroom. The BJR is intended, to some extent, to address the reality that good business decisions require creativity. It gives directors room to be creative in their decision-making through lower risk. As discussed above, AI systems are not capable of developing a genuinely novel idea, because systems that are "trained on historical data ... will be "backward-looking".¹⁰⁰ Because of their statistical nature, AI systems are risk averse. They will make business-decision recommendations only on the basis of what they 'know' to have been previously successful. A lack of creativity in decision-making leads to a homogeneity of decisions. An inability to evolve the role of the corporation within society is also a concern. The duty of care is not intended to support variations on the same decision endlessly, hence the importance of retaining the BJR.

(e) The Presence of Artificial Intelligence

A limitation of the analysis in the foregoing scenarios is the use of ChatGPT as the AI system that Director A consults. There are many tools that could constitute AI which Director A might be more readily expected to use than ChatGPT. There are more specialized AI tools, like the Salesforce 'Einstein' program alluded to earlier. The use of ChatGPT gives specificity to the analysis here, since many discussions in this area do not clearly envision or articulate what is meant by "AI." In some analyses, an AI system is conceptualized as some kind of monolithic robot or algorithm.

Directors may also use tools which include AI programs which are not immediately disclosed, like upgraded Microsoft software. The presence of an AI system is not always apparent and a lack of awareness of the presence of AI can be problematic. During Standing Committee testimony on Canada's AI legislation, Bianca Wylie observed: "How would you know if you have been harmed by artificial intelligence? It is ... like asbestos. It's in things where you don't even know it's there."¹⁰¹

*222 While the analogy was drawn in the context of discussing 'harms' deriving from AI, the principle extends to use of the technology. Directors may unwittingly use an AI-powered system, overly rely on it, and, as a result, breach their duty of care.

Outside the use of what is clearly an AI protocol -- like ChatGPT -- directors should also be cautious and inform themselves about where AI might be encountered in their decision-making process and how the system should be appropriately used. A failure to fully scrutinize the tools by which the decision is being made could result in a breach of the duty of care through overreliance. Directors must retain authority in decision-making.

(f) Duty to Use?

These scenarios have not discussed instances where AI can be helpful to directors, and in fact improve decision making. As Mertens notes "AI is able to complement the broad capabilities and knowledge of the human board members by providing them with a clear analysis of intangible mountains of data, and therefore increase the pace at which difficult decisions are taken."¹⁰² Certainly, there are instances where the use of cutting-edge technologies is appropriate -- and indeed required - in business activities. For example, no one expects a business owner to be maintaining a physical ledger. An excel spreadsheet, or other similar software, is the expected method of tracking cash flows in modern business. AI may take on such a role in some applications in businesses. An appropriately tailored and understood AI tool might indeed significantly aid directorial decision-making and contribute to the director meeting the duty of care.

Because of the data processing capabilities and other functionalities that AI presents for businesses, some commentators have suggested that in order to meet their duty of care, directors should be required to consult or utilize AI systems. As Zhao explains:

[I]f relying on AI could lead to more informed decisions, utilizing its capabilities to process information and make predictions using large data sets, it is then reasonable to expect directors to consult AI when making strategic decisions ... Directors are already bound to undertake due diligence before making decisions to ensure they have adequate information; otherwise, they may violate their duty of care."¹⁰³

The analysis in the scenarios makes clear that to require a director to use AI in order to meet their duty of care would not be appropriate. Such a requirement ***223** could lead to overreliance on the system, and the AI assuming a problematically authoritative role. Further, a requirement for directors to use AI in decision-making is largely predicated on the idea that AI makes 'better' decisions because it is able to process more information and "mitigate the human lack of capability to understand complex data and subsequently choose between the options available."¹⁰⁴ In this view, AI systems are more rational than human decision-makers, again revealing a conception of AI as 'neutral.'

This 'all or nothing' theory is also problematic because the AI systems it references are amorphous. The AI systems proposed for use in these discussions are not clearly defined or discussed, beyond broad principles. There may be instances where, to become fully informed in the face of a huge amounts of data, a director is effectively required to use AI analytics or tools to have done their due diligence and meet the duty of care.¹⁰⁵ But the use of AI in such a case would not result in an overreliance on the tool. In such a use, the AI system would still be in the role of assistant. The analysis would fall within Scenario 2. Characterizing AI as 'more rational' runs the risk of directors relying on the authoritative AI illustrated in Scenario 3, hence resulting in a failure to meet the duty of care.

A requirement upon directors to use AI in order to meet the duty of care is also unsupported by the BJR. To shelter under the BJR, a decision has to be reasonable. Reasonability alone is sufficient. Under the duty of care "directors do not have the duty to gather every piece of information available, or to maximize the informed basis on which they take their decisions."¹⁰⁶ A mandatory AI-system consultation would effectively create a such a requirement. This would overburden the duty of care, making it onerous. To have met the duty of care under this conception, the director would have to have a much more granular understanding of the data -- or have used a machine to give them such an understanding -- in order to have met their duty of care. This would make it challenging for directors to arrive at any decision without opening themselves to liability, since a director rarely knows the totality of the information relevant to making a given decision.

IV. CONCLUSION

In sum, the director's duty of care under the *CBCA* has ramifications for the directorial use of AI tools in decision-making. This becomes clear with an *224 understanding of the technical operation of AI systems. These are systems that rely predominantly on statistical analyses, predictive goals, and training data. A consideration of the modes of reasoning that AI systems employ -- and cannot employ -- is also illuminating. The examination reveals that AI systems in their current form cannot be "reasonable" in the legal sense of the word. Although the output from these systems may make them appear capable of decision-making, AIs do not have the capacity to make judgments. They are systems that have no ability to perform abductive reasoning patterns, as well as failing to engage with reasoning in a contextual and subjective manner.

Examination of three scenarios of directorial use of ChatGPT suggests the impact that this lack of reasonableness has on the duty of care, particularly with respect to the extent to which it is permissible for directors to delegate reasoning to the machine. The directorial duty of care is likely to be met in instances where directors use AI tools to assist their thinking in arriving at a decision. However, total reliance on the output of a tool in decision-making is not appropriate. Such a reliance undercuts the duty of care, first by absolving directors of almost all responsibility in respect of decision making, and second by removing reasonability from the equation. Directors should take care to rely on AI tools only as assistants and avoid relying on the systems as authorities. Such awareness is particularly important as AI becomes increasingly prevalent, and less clearly present, in business tools and systems.

*225 BIBLIOGRAPHY

Primary Sources

Business Corporations Act, RSO 1190, c B.16.

Canada Business Corporations Act, RSC 1985, c C-44.

BCE Inc., Re, 2008 SCC 69, 2008 CarswellQue 12595, 2008 CarswellQue 12596 (S.C.C.).

People's Department Stores Ltd. (1992) Inc., Re, 2004 CarswellQue 2862, 2004 CarswellQue 2863, [2004] 3 S.C.R. 461 (S.C.C.).

City Equitable Fire Insurance Co., Re (1924), [1925] 1 Ch. 407 (Eng. Ch. Div.), affirmed [1924] 1 Ch. 407 at 500 (C.A.).

Secondary Sources

Belcastro, Thomas, "Getting on Board with Robots: How the Business Judgement Rule Should Apply to Artificially Intelligent Devices Serving as Members of a Corporate Board" (2019) 4:1 Georgetown L T Rev 263.

Bender, Emily M et al, "On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?" (Paper delivered at ACM Conference on Fairness, Accountability and Transparency, 1 March 2021) online: < doi.org/10.1145/3442188.3445922>.

Berkemer, Rainer & Markus Grotte, "Learning Algorithms: What is Artificial Intelligence Really Capable Of?" in Peter Klimczak & Christer Petersen, eds, *AI: Limits and Prospects of Artificial Intelligence*, (Bielefeld: transcript Verlag, 2023) 9.

Brownell, Briana, "Ted-Ed: How does artificial intelligence learn?" (11 March 2021) online (video): [https://perma.cc/SZG3-TGS5].

Burridge, Nicky, "Artificial intelligence gets a seat in the boardroom", *Nikkei Asia* (10 May 2017), online: [https://perma.cc/7ZCV-5FK2].

Camp, Sophie, "Why everyone in the boardroom needs AI" online: [https://perma.cc/NY3V-DAC4].

Cantwell Smith, Briana, *The Promise of Artificial Intelligence: Reckoning and Judgement* (Cambridge: MIT Press, 2019).

Cuthbertson, Anthony, "Company that made an AI its chief executive sees stocks climb", *The Independent* (16 March 2023), online: [<https://perma.cc/8BN4-KZG6>].

Delacroix, Sylvie, "Automated Systems and the Need for Change" in Simon Deakin and Christopher Markou, eds, *Is Law Computable? Critical Perspectives on Law and Artificial Intelligence*, (Oxford: Hart, 2020) 184.

Deloitte, *Canada's AI imperative: Start, Scale, Succeed*, (Deloitte Canada, 2023).

Handbook on Law and Technology, (Cheltenham: Edward Elgar Publishing, 2023) 427.

Hannigan, Tim, Ian P. McCarthy & Andre Spicer "Beware of Botshit: How to Manage the Epistemic Risks of Generative Chatbots" (2024) 67:5 *Bus Horizons* 471-486.

*226 Heikkinen, Tatjana et al., "Artificial intelligence and the law: can we and should we regulate AI systems?" in Bartosz Brozek, Olesia Kanevskaia, and Przemyslaw Palka, eds, *Research*

Hildebrandt, Mireille, "Code-driven Law: Freezing the Future and Scaling the Past" in Simon Deakin and Christopher Markou, eds, *Is Law Computable? Critical Perspectives on Law and Artificial Intelligence*, (Oxford: Hart, 2020) 90.

House of Commons, Standing Committee on Industry and Technology, *Evidence*, 44-1, No 102 (7 December 2023) (Bianca Wylie).

Larson, Erik J., *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021).

Larson, Erik, "If AI's Don't Know What They're Doing, Can We Hope to Explain It?" (23 February 2024), online: [<https://perma.cc/5MRZ-8RKP>].

Larson, Erik, "Why, Despite All the Hype We Hear, AI is Not "One Of Us" (22 February 2024), online: [<https://perma.cc/CN9Z-NALZ>].

Lehr, David & Paul Ohm, "Playing with the Data: What Legal Scholars Should Learn About Machine Learning" (2017) 51 *UCDL Rev* 653.

Mertens, Floris, "The use of artificial intelligence in corporate decision-making at board level: a preliminary legal analysis" (2023), University of Ghent Financial Law Institute, Working Paper No 2023-01.

Mok, Aaron, "Google DeepMind cofounder says AI can act like an entrepreneur and inventor in the next five years" (18 January 2024) online: [<https://perma.cc/2SEC-8PNN>].

Moslein, Florian, "Robots in the Boardroom: Artificial Intelligence and Corporate Law" in Woodrow Barfield and Ugo Pagallo, eds, *Research Handbook on the Law of Artificial Intelligence*, (Cheltenham: Edward Elgar Publishing, 2018) 649.

Norvig, Peter & Stuart Russell, *Artificial Intelligence: A Modern Approach*, (Hoboken: Pearson Publishing, 2021).

Petrin, Martin, "Corporate Management in the Age of AI" (2019) *Colum Bus. L. Rev* 2019:3 965.

Picciau, Chiara, "The (Un)Predictable Impact of Technology on Corporate Governance" (2021) 17:1 *Hastings Bus L J* 67.

Purtill, Corinne, "The emails that brought down Enron still shape our daily lives" (15 February 2019), online: [<https://perma.cc/28WB-ZP9P>]. Rajendran, Janarthanan, "Building AI Systems: Training Data, Machine Learning & Algorithms" (Presentation delivered at LAWS 2019 Class, March 4, 2024) [unpublished].

Reiter, Barry J., “Doing Your Job as a Director” in Barry J Reiter, Gary S.A. Solway & Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49.

Robertson, Ian, “Soon, artificial intelligence will be running companies -- rise of the robo-director”, *The Globe and Mail* (18 May 2023), online: [<https://perma.cc/3A28-KH6Z>].

*227 Smith, Justin, *The Internet Is Not What You Think It Is: A History, a Philosophy, a Warning*, (Hoboken: Pearson Publishing, 2021).

Welling, Bruce, *Corporate Law in Canada: The Governing Principles*, 3rd ed (Mudgeeraba, Queensland: Scribblers, 2006).

Zerilli, John Z et al, *A Citizen's Guide to Artificial Intelligence*, (Cambridge: MIT Press, 2021).

Zhao, Jingchen, “Artificial Intelligence and Corporate Decisions: Fantasy, Reality or Destiny?” (2022) 17:4 Catholic U L Rev 663.

Footnotes

a1 Margaret Wilson received her Juris Doctor with a Certificate in Law and Technology from the Schulich School of Law at Dalhousie University in 2024. She also holds graduate and undergraduate degrees in materials engineering. Margaret is thankful for the support and feedback of Professor Liam McHugh-Russell in the development of this article.

1 Stuart Russell & Peter Norvig, *Artificial Intelligence: A Modern Approach*, (Hoboken: Pearson Publishing, 2021) at 2.

2 Aaron Mok, “Google DeepMind cofounder says AI can act like an entrepreneur and inventor in the next five years” (18 January 2024) online: < <https://www.businessinsider.com/google-deepmind-ai-could-run-businesses-in-five-years-20241#:~:text=AI%20leader%20Mustafa%20Suleyman%20predicts,the%202024%20World%20Economic%20Forum> > [<https://perma.cc/2SEC-8PNN>].

3 Nicky Burrige, “Artificial intelligence gets a seat in the boardroom”, *Nikkei Asia* (10 May 2017), online: < <https://asia.nikkei.com/Business/Artificial-intelligence-gets-a-seat-in-the-boardroom> > [<https://perma.cc/7ZCV-5FK2>].

4 Anthony Cuthbertson, “Company that made an AI its chief executive sees stocks climb”, *The Independent* (16 March 2023), online: < <https://www.independent.co.uk/tech/ai-ceo-artificial-intelligence-b2302091.html> > [<https://perma.cc/8BN4-KZG6>].

5 Sophie Camp, “Why everyone in the boardroom needs AI” (last visited 1 April 2024) online: < <https://outsideinsight.com/insights/why-everyone-in-the-boardroom-needs-ai/> > [<https://perma.cc/NY3V-DAC4>].

6 Sophie Camp, “Why everyone in the boardroom needs AI” (last visited 1 April 2024) online: < <https://outsideinsight.com/insights/why-everyone-in-the-boardroom-needs-ai/> > [<https://perma.cc/NY3V-DAC4>].

7 *Canada Business Corporations Act*, R.S.C. 1985, c. C-44, ss105(1) [CBCA].

8 See for example, *Business Corporations Act*, R.S.O. 1190, c. B.16, ss 118(1).

9 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 49.

10 *City Equitable Fire Insurance Co., Re* (1924), [1925] 1 Ch. 407 (Eng. Ch. Div.), affirmed [1924] 1 Ch. 407 at 500 (C.A.).

11 Bruce Welling, *Corporate Law in Canada: The Governing Principles*, 3rd ed (Mudgeeraba, Queensland: Scribblers, 2006) at 325.

12 *Canada Business Corporations Act*, R.S.C. 1985, c. C-44, ss 122(1)(b).

- 13 [People's Department Stores Ltd. \(1992\) Inc., Re](#), 2004 CarswellQue 2862, 2004 CarswellQue 2863, [2004] 3 S.C.R. 461 (S.C.C.) at para. 63 [Peoples].
- 14 Bruce Welling, *Corporate Law in Canada: The Governing Principles*, 3rd ed (Mudgeeraba, Queensland: Scribblers, 2006) at 329.
- 15 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 67.
- 16 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 71.
- 17 Bruce Welling, *Corporate Law in Canada: The Governing Principles*, 3rd ed (Mudgeeraba, Queensland: Scribblers, 2006) at 328.
- 18 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 52.
- 19 [People's Department Stores Ltd. \(1992\) Inc., Re](#), 2004 CarswellQue 2862, 2004 CarswellQue 2863, [2004] 3 S.C.R. 461 (S.C.C.) at para 57; [BCE Inc., Re](#), 2008 SCC 69, 2008 CarswellQue 12595, 2008 CarswellQue 12596 (S.C.C.) at para. 104.
- 20 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 73.
- 21 Bruce Welling, *Corporate Law in Canada: The Governing Principles*, 3rd ed (Mudgeeraba, Queensland: Scribblers, 2006) at 326.
- 22 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 67.
- 23 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 67.
- 24 [People's Department Stores Ltd. \(1992\) Inc., Re](#), 2004 CarswellQue 2862, 2004 CarswellQue 2863, [2004] 3 S.C.R. 461 (S.C.C.) at para 67.
- 25 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 71.
- 26 Barry J. Reiter, Doing Your Job as a Director in Barry J Reiter, Gary S.A. Solway, and Jesslyn G Maurier, eds, 7th ed, *Directors' Duties in Canada*, (Toronto: LexisNexis, 2021) 49 at 71.
- 27 Bruce Welling, *Corporate Law in Canada: The Governing Principles*, 3rd ed (Mudgeeraba, Queensland: Scribblers, 2006) at 329-30.
- 28 *Canada Business Corporations Act*, R.S.C. 1985, c. C-44, ss 123(4)(b).
- 29 Bruce Welling, *Corporate Law in Canada: The Governing Principles*, 3rd ed (Mudgeeraba, Queensland: Scribblers, 2006) at 330-1.
- 30 John Zerilli et al, *A Citizen's Guide to Artificial Intelligence*, (Cambridge: MIT Press, 2021) at 1.
- 31 Riika Koulu et al, "Artificial intelligence and the law: can we and should we regulate AI systems?" in Bartosz Brożek, Olia Kanevskaia, & Przemysław Palka, eds, *Research Handbook on Law and Technology*, (Cheltenham: Edward Elgar Publishing, 2023) 427 at 428.

- 32 Riika Koulu et al, “Artificial intelligence and the law: can we and should we regulate AI systems?” in Bartosz Brożek, Oľia Kanevskaia, & Przemysław Palka, eds, *Research Handbook on Law and Technology*, (Cheltenham: Edward Elgar Publishing, 2023) 427 at 433.
- 33 John Zerilli et al, *A Citizen's Guide to Artificial Intelligence*, (Cambridge: MIT Press, 2021) at 4.
- 34 Stuart Russell & Peter Norvig, *Artificial Intelligence: A Modern Approach*, (Hoboken: Pearson Publishing, 2021) at 651.
- 35 David Lehr & Paul Ohm, “Playing with the Data: What Legal Scholars Should Learn About Machine Learning” (2017) 51 UCDL Rev 653-717 at 674.
- 36 David Lehr & Paul Ohm, “Playing with the Data: What Legal Scholars Should Learn About Machine Learning” (2017) 51 UCDL Rev 653-717 at 674.
- 37 David Lehr & Paul Ohm, “Playing with the Data: What Legal Scholars Should Learn About Machine Learning” (2017) 51 UCDL Rev 653-717 at 672-3.
- 38 John Zerilli et al, *A Citizen's Guide to Artificial Intelligence*, (Cambridge: MIT Press, 2021) at 4.
- 39 Corinne Purtill, “The emails that brought down Enron still shape our daily lives” (15 February 2019), online: < <https://qz.com/work/1546565/the-emails-that-brought-down-enron-still-shape-our-daily-lives/> [https://perma.cc/28WB-ZP9P].
- 40 Corinne Purtill, “The emails that brought down Enron still shape our daily lives” (15 February 2019), online: < <https://qz.com/work/1546565/the-emails-that-brought-down-enron-still-shape-our-daily-lives/> [https://perma.cc/28WB-ZP9P].
- 41 Stuart Russell & Peter Norvig, *Artificial Intelligence: A Modern Approach*, (Hoboken: Pearson Publishing, 2021) at 676-7.
- 42 John Zerilli et al, *A Citizen's Guide to Artificial Intelligence*, (Cambridge: MIT Press, 2021) at 11.
- 43 Erik Larson, “Why, Despite All the Hype We Hear, AI is Not “One Of Us”D'D' (22cFebruary 2024), online: < <https://mindmatters.ai/2024/02/why-despite-all-the-hype-we-hear-ai-is-not-one-of-us/> [https://perma.cc/CN9Z-NALZ].
- 44 John Zerilli et al, *A Citizen's Guide to Artificial Intelligence*, (Cambridge: MIT Press, 2021) at 12.
- 45 Stuart Russell & Peter Norvig, *Artificial Intelligence: A Modern Approach*, (Hoboken: Pearson Publishing, 2021) at 751.
- 46 John Zerilli et al, *A Citizen's Guide to Artificial Intelligence*, (Cambridge: MIT Press, 2021) at 14.
- 47 Briana Brownell, “Ted-Ed: How does artificial intelligence learn?” (11 March 2021) at 0h0m44s, online (video): < <https://www.youtube.com/watch?v=0yCJMt9Mx9c> [https://perma.cc/SZG3-TGS5].
- 48 David Lehr & Paul Ohm, “Playing with the Data: What Legal Scholars Should Learn About Machine Learning” (2017) 51 UCDL Rev 653-717 at 673.
- 49 Stuart Russell & Peter Norvig, *Artificial Intelligence: A Modern Approach*, (Hoboken: Pearson Publishing, 2021) at 653.
- 50 Stuart Russell & Peter Norvig, *Artificial Intelligence: A Modern Approach*, (Hoboken: Pearson Publishing, 2021) at 653.
- 51 John Zerilli et al, *A Citizen's Guide to Artificial Intelligence*, (Cambridge: MIT Press, 2021) at 14.
- 52 Janarthanan Rajendran, “Building AI Systems: Training Data, Machine Learning & Algorithms” (Presentation delivered at LAWS 2019 Class, March 4, 2024) [unpublished].
- 53 Tim Hannigan, Ian P. McCarthy, & Andre Spicer, “Beware of Botshit: How to Manage the Epistemic Risks of Generative Chatbots” (2024) 67:5 Bus Horizons at 472.

- 54 Tim Hannigan, Ian P. McCarthy, & Andre Spicer, “Beware of Botshit: How to Manage the Epistemic Risks of Generative Chatbots” (2024) 67:5 Bus Horizons at 472.
- 55 Justin Smith, *The Internet Is Not What You Think It Is: A History, a Philosophy, a Warning*, (Hoboken: Pearson Publishing, 2021) at 45-6.
- 56 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 4.
- 57 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 110.
- 58 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 189.
- 59 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 115.
- 60 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 189-90.
- 61 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 176.
- 62 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 161.
- 63 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 4.
- 64 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 166.
- 65 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 163.
- 66 Erik J Larson, *The Myth of Artificial Intelligence: Why Computers Can't Think the Way We Do*, (Cambridge, Massachusetts: The Belknap Press of Harvard University, 2021) at 183.
- 67 Erik Larson, “If AI's Don't Know What They're Doing, Can We Hope to Explain It?” (23 February 2024), online: < <https://mindmatters.ai/2024/02/if-ais-dont-know-what-theyre-doing-can-we-hope-to-explain-it/> [https://perma.cc/5MRZ-8RKP].
- 68 Erik Larson, “If AI's Don't Know What They're Doing, Can We Hope to Explain It?” (23 February 2024), online: < <https://mindmatters.ai/2024/02/if-ais-dont-know-what-theyre-doing-can-we-hope-to-explain-it/> [https://perma.cc/5MRZ-8RKP].at 174-5.
- 69 Erik Larson, “If AI's Don't Know What They're Doing, Can We Hope to Explain It?” (23 February 2024), online: < <https://mindmatters.ai/2024/02/if-ais-dont-know-what-theyre-doing-can-we-hope-to-explain-it/> [https://perma.cc/5MRZ-8RKP] at 189-90.
- 70 Erik Larson, “If AI's Don't Know What They're Doing, Can We Hope to Explain It?” (23 February 2024), online: < <https://mindmatters.ai/2024/02/if-ais-dont-know-what-theyre-doing-can-we-hope-to-explain-it/> [https://perma.cc/5MRZ-8RKP] at 180.
- 71 Brian Cantwell Smith, *The Promise of Artificial Intelligence: Reckoning and Judgement* (Cambridge: MIT Press, 2019) at 80.

- 72 Rainer Berkemer & Markus Grotte, “Learning Algorithms: What is Artificial Intelligence Really Capable Of?” in Peter Klimcazk & Christer Petersen, eds, *AI: Limits and Prospects of Artificial Intelligence*, (Bielefeld: transcript Verlag, 2023) 9 at 34.
- 73 Emily M Bender et al, “On the Dangers of Stochastic Parrots: Can Language Models Be Too Big?” (Paper delivered at ACM Conference on Fairness, Accountability and Transparency, 1 March 2021) at 616 online: doi.org/10.1145/3442188.3445922.
- 74 Brian Cantwell Smith, *The Promise of Artificial Intelligence: Reckoning and Judgement* (Cambridge: MIT Press, 2019) at 80 at 109-110.
- 75 Brian Cantwell Smith, *The Promise of Artificial Intelligence: Reckoning and Judgement* (Cambridge: MIT Press, 2019) at 80 at 108.
- 76 Rainer Berkemer & Markus Grotte, “Learning Algorithms: What is Artificial Intelligence Really Capable Of?” in Peter Klimcazk & Christer Petersen, eds, *AI: Limits and Prospects of Artificial Intelligence*, (Bielefeld: transcript Verlag, 2023) 9 at 34.
- 77 Stuart Russell & Peter Norvig, *Artificial Intelligence: A Modern Approach*, (Hoboken: Pearson Publishing, 2021) at 986.
- 78 Rainer Berkemer & Markus Grotte, “Learning Algorithms: What is Artificial Intelligence Really Capable Of?” in Peter Klimcazk & Christer Petersen, eds, *AI: Limits and Prospects of Artificial Intelligence*, (Bielefeld: transcript Verlag, 2023) 9 at 37.
- 79 Rainer Berkemer & Markus Grotte, “Learning Algorithms: What is Artificial Intelligence Really Capable Of?” in Peter Klimcazk & Christer Petersen, eds, *AI: Limits and Prospects of Artificial Intelligence*, (Bielefeld: transcript Verlag, 2023) 9 at 35-6.
- 80 Brian Cantwell Smith, *The Promise of Artificial Intelligence: Reckoning and Judgement* (Cambridge: MIT Press, 2019) at 80 at xvii.
- 81 Brian Cantwell Smith, *The Promise of Artificial Intelligence: Reckoning and Judgement* (Cambridge: MIT Press, 2019) at 80 at xv-xvi.
- 82 Rainer Berkemer & Markus Grotte, “Learning Algorithms: What is Artificial Intelligence Really Capable Of?” in Peter Klimcazk & Christer Petersen, eds, *AI: Limits and Prospects of Artificial Intelligence*, (Bielefeld: transcript Verlag, 2023) 9 at 35.
- 83 Brian Cantwell Smith, *The Promise of Artificial Intelligence: Reckoning and Judgement* (Cambridge: MIT Press, 2019) at 80 at xix.
- 84 Deloitte, *Canada's AI imperative: Start, Scale, Succeed*, (Deloitte Canada, 2023) at 4.
- 85 Ian Robertson, “Soon, artificial intelligence will be running companies -- rise of the robo-director”, *The Globe and Mail* (18 May 2023), online: < <https://www.theglobeandmail.com/business/commentary/article-artificial-intelligence-robo-director-companies/> [https://perma.cc/3A28-KH6Z].
- 86 Florian Moslein, “Robots in the Boardroom: Artificial Intelligence and Corporate Law” in Woodrow Barfield and Ugo Pagallo, eds, *Research Handbook on the Law of Artificial Intelligence*, (Cheltenham: Edward Elgar Publishing, 2018) 649 at 657.
- 87 Florian Moslein, “Robots in the Boardroom: Artificial Intelligence and Corporate Law” in Woodrow Barfield and Ugo Pagallo, eds, *Research Handbook on the Law of Artificial Intelligence*, (Cheltenham: Edward Elgar Publishing, 2018) 649 at 657.
- 88 *Canada Business Corporations Act*, R.S.C. 1985, c. C-44, ss 123(4)-123(5).
- 89 Florian Moslein, “Robots in the Boardroom: Artificial Intelligence and Corporate Law” in Woodrow Barfield and Ugo Pagallo, eds, *Research Handbook on the Law of Artificial Intelligence*, (Cheltenham: Edward Elgar Publishing, 2018) 649 at 658.
- 90 Florian Moslein, “Robots in the Boardroom: Artificial Intelligence and Corporate Law” in Woodrow Barfield and Ugo Pagallo, eds, *Research Handbook on the Law of Artificial Intelligence*, (Cheltenham: Edward Elgar Publishing, 2018) 649.
- 91 Sophie Camp, “Why everyone in the boardroom needs AI” (last visited 1 April 2024) online: < <https://outsideinsight.com/insights/why-everyone-in-the-boardroom-needsai/> [https://perma.cc/NY3V-DAC4].

- 92 Florian Moslein, “Robots in the Boardroom: Artificial Intelligence and Corporate Law” in Woodrow Barfield and Ugo Pagallo, eds, *Research Handbook on the Law of Artificial Intelligence*, (Cheltenham: Edward Elgar Publishing, 2018) 649 at 663.
- 93 Martin Petrin, “Corporate Management in the Age of AI” (2019) *Colum Bus. L. Rev* 2019:3 965 at 984.
- 94 Martin Petrin, “Corporate Management in the Age of AI” (2019) *Colum Bus. L. Rev* 2019:3 965 at 984.
- 95 Mireille Hildebrandt, “Code-driven Law: Freezing the Future and Scaling the Past” in Simon Deakin and Christopher Markou, eds, *Is Law Computable? Critical Perspectives on Law and Artificial Intelligence*, (Oxford: Hart, 2020) 90 at 91.
- 96 Martin Petrin, “Corporate Management in the Age of AI” (2019) *Colum Bus. L. Rev* 2019:3 965 at 984 at 657.
- 97 Chiara Picciau, “The (Un)Predictable Impact of Technology on Corporate Governance” (2021) 17:1 *Hastings Bus L J* 67 at 132.
- 98 *People's Department Stores Ltd. (1992) Inc., Re*, 2004 CarswellQue 2862, 2004 CarswellQue 2863, [2004] 3 S.C.R. 461 (S.C.C.) at para 67.
- 99 Thomas Belcastro, “Getting on Board with Robots: How the Business Judgement Rule Should Apply to Artificially Intelligent Devices Serving as Members of a Corporate Board” (2019) 4:1 *Georgetown L T Rev* 263 at 276.
- 100 Sylvie Delacroix, “Automated Systems and the Need for Change” in Simon Deakin and Christopher Markou, eds, *Is Law Computable? Critical Perspectives on Law and Artificial Intelligence*, (Oxford: Hart, 2020) 184 at 185.
- 101 House of Commons, Standing Committee on Industry and Technology, *Evidence*, 44-1, No 102 (7 December 2023) at 16:50 (Bianca Wylie).
- 102 Floris Mertens, “The use of artificial intelligence in corporate decision-making at board level: a preliminary legal analysis” (2023), University of Ghent Financial Law Institute, Working Paper No 2023-01 at 9.
- 103 Jingchen Zhao, “Artificial Intelligence and Corporate Decisions: Fantasy, Reality or Destiny” (2022) 17:4 *Catholic U L Rev* 663 at 686.
- 104 Jingchen Zhao, “Artificial Intelligence and Corporate Decisions: Fantasy, Reality or Destiny” (2022) 17:4 *Catholic U L Rev* 663 at 675.
- 105 Floris Mertens, “The use of artificial intelligence in corporate decision-making at board level: a preliminary legal analysis” (2023), University of Ghent Financial Law Institute, Working Paper No 2023-01 at 21.
- 106 Florian Moslein, “Robots in the Boardroom: Artificial Intelligence and Corporate Law” in Woodrow Barfield and Ugo Pagallo, eds, *Research Handbook on the Law of Artificial Intelligence*, (Cheltenham: Edward Elgar Publishing, 2018) 649 at 661.

22 CANJLT 201