

DOI: https://doi.org/10.48009/4_iis_2024_116

AI and management: navigating the alignment problem for ethical and effective decision-making

Robert Boncella, *Washburn University*, bob.boncella@washburn.edu

Abstract

In tutorial format, this paper explores the intricate relationship between Artificial Intelligence (AI) and managerial decision-making, emphasizing the alignment problem—a critical challenge in ensuring AI systems align with human values and ethical standards. It defines AI and examines various methods, including machine learning algorithms, natural language processing, recommendation systems, sentiment analysis, data visualization, anomaly detection, expert systems, neural networks, and deep learning. The alignment problem is dissected into key aspects such as defining goals, value misalignment, robustness and safety, long-term consequences, and interpreting human preferences. The types of AI most affected by the alignment problem, including deep learning, reinforcement learning, generative adversarial networks, evolutionary algorithms, and open-ended learning systems, are highlighted. Practical examples illustrate how biases and misalignments manifest in business processes like recruitment, credit scoring, and automated trading. The paper also categorizes managerial decisions and discusses how AI methods support these decisions while addressing alignment issues. This document underscores the importance of aligning AI systems with human values to ensure ethical and effective outcomes in business environments. This paper is intended to be a road map of possible research topics on the effect of AI on Business Decision Making through the lens of the Alignment Problem.

Keywords: artificial intelligence (AI), alignment problem, managerial decision-making, machine learning, ethics, business process

Introduction

Artificial Intelligence (AI) has become pivotal in modern technological advancements, driving significant transformations across various sectors. This article delves into the intricate relationship between AI and managerial decision-making, emphasizing the alignment problem—a critical challenge in ensuring AI systems work harmoniously with human values and ethical standards. By simulating human intelligence, AI systems are designed to perform tasks that typically require human cognition, such as learning, problem-solving, and decision-making.

The paper begins by defining AI and exploring various AI methods, including machine learning algorithms, natural language processing, recommendation systems, sentiment analysis, data visualization, anomaly detection, expert systems, neural networks, and deep learning. Each method's capabilities and applications are outlined, displaying their roles in enhancing business processes and decision-making. A significant focus is placed on the alignment problem in AI, which addresses the challenge of ensuring AI systems act in ways consistent with human values and intentions. Key aspects of this problem are discussed, including defining goals, value misalignment, robustness and safety, long-term consequences, and interpreting human preferences. The types of AI most affected by the alignment problem, such as deep learning, reinforcement learning, generative adversarial networks, evolutionary algorithms, and open-ended learning systems, are examined to highlight the potential risks and ethical considerations.

Furthermore, the document gives practical examples of the alignment problem in AI-supported business processes, illustrating how biases and misalignments can manifest in areas like recruitment, credit scoring, automated trading, customer service, supply chain optimization, predictive policing, healthcare, personalized marketing, employee monitoring, and loan approval.

Finally, managerial decision-making is explored, categorizing decisions into operational, managerial, tactical, strategic, ad-hoc analysis, collaborative, knowledge-driven, and data-driven types. The paper discusses how AI methods support these decision types and addresses the associated alignment issues, providing a comprehensive overview of AI's impact on business decision-making. This article offers insights into the complex interplay between AI technologies and managerial decision-making, underscoring the importance of aligning AI systems with human values to ensure ethical and effective outcomes in business environments.

Artificial intelligence and the alignment problem

Artificial Intelligence (AI) - refers to the simulation of human intelligence by machines that are programmed to think like humans and mimic their actions. AI is shown by any machine with traits associated with a human mind, such as learning and problem-solving. (Duan, 2023) A machine is any computing system used to run AI algorithms and applications, enabling it to perform tasks that require intelligence. The defining characteristic of these machines is their ability to process data and execute algorithms to perform intelligent tasks. Achieving this functionality involves various forms of machine learning, in which the machine is trained on a dataset to make predictions or decisions without being explicitly programmed to perform the task. This training (learning) can take two forms: Supervised or Unsupervised.

Supervised learning is a type of machine learning where the algorithm is trained on a labeled dataset. This means that for each input in the training set, there is a corresponding output (label). The algorithm learns to map inputs to outputs by finding patterns in the data. In supervised learning, the goal is to learn a function that maps an input to an output based on example input-output pairs. It involves two main processes: training and testing. During the training phase, the algorithm makes predictions on the training data and adjusts its parameters based on the error of its predictions. This process continues until the algorithm achieves a desired level of accuracy. Once trained, the model can be used to make predictions on new, unseen data.

Unsupervised learning is a type of machine learning where the algorithm is trained on unlabeled data. The system tries to learn the underlying patterns and structure from the input data without any explicit instructions on what to predict. In unsupervised learning, the algorithm attempts to infer the natural structure present within a set of data points. Unlike supervised learning, there are no labels or predefined outcomes. The most common tasks in unsupervised learning are clustering and association. Clustering involves grouping the data into clusters of similar items, while association involves discovering interesting relationships between variables in large datasets.

Sample of artificial intelligence methods.

Table 1 shows the AI methods that will be used in this paper. This is not an exhaustive list; the interested reader should consult it. (Russell & Norvig 2020)

Table 1 Types of AI Methods

Type of Method	Definition	Example
Machine Learning Algorithms – (ML)	Algorithms that enable machines to process data, learn from it, and then use what they have learned to make informed decisions. (Russell & Norvig 2020)	Predictive analytics in financial forecasting.
Natural Language Processing - (NLP)	Techniques that enable machines to understand, interpret, and generate human language meaningfully. (Russell & Norvig 2020)	chatbots like Siri, Alexa, and Google Assistant.
Recommendation Systems:	Systems that filter and predict user preferences based on past behavior, preferences, and similar user profiles to suggest relevant items or content. (Ricci et al.,2011)	An online streaming service suggesting movies based on a user's earlier viewing history and ratings.
Sentiment Analysis:	Computationally naming and categorizing opinions expressed in a text to determine whether the writer's attitude is positive, negative, or neutral. (Liu, B. 2012)	Analyzing social media posts to gauge public sentiment about a new product launch.
Data Visualization:	The graphical representation of information and data enables easy understanding and interpretation of complex data. (Tuft, E. R. 2001)	Using a dashboard to represent sales data trends over time visually.
Anomaly Detection:	Naming rare items, events, or outlier observations that raise suspicions by differing significantly from most of the data. (Chandola et al., 2009).	Detecting fraudulent transactions in banking by identifying patterns that deviate from typical user behavior.
Expert Systems	Computer systems that emulate the decision-making abilities of a human expert in particular domains. (Russell & Norvig 2020)	Medical diagnosis systems.
Neural Networks	Algorithms designed to recognize patterns and interpret sensory data through machine perception, labeling, and clustering of raw input. (Russell & Norvig 2020)	include handwriting recognition or image classification systems.
Deep Learning	A subset of machine learning that uses neural networks with many layers (hence "deep" learning) to analyze several factors of data. (Russell & Norvig 2020)	Voice recognition systems like Google's Voice Search or ChatGPT.

The alignment problem (TAP) in AI is the problem of ensuring that Artificial Intelligence (AIs) and artificial intelligence agents (AIAs) act in ways that are aligned with human values, ethics, and intentions. The problem can be broken down into several components. The key aspects of the alignment problem are listed below (Christian, 2020), (Russell, 2019), (Amodei et al., 2016), and (Bostrom, (2014).

Defining Objectives: Defining what we want an AI system to do can be complex. A misalignment between human intent and the goal set for the AI can lead to unintended behavior.

Value Misalignment: Human values can be complex, multifaceted, and sometimes inconsistent. Embedding these values into an AI system that accurately reflects human ethics and morals is a substantial challenge.

Robustness and Safety: Ensuring that AI behaves safely under various conditions, including those not met during training, is a significant concern.

Long-term Consequences: The alignment problem also considers AI's long-term impact and development trajectory. If AI evolves in a way not aligned with human interests, it could have potentially harmful consequences.

Interpreting Human Preferences: Understanding and analyzing human preferences and instructions can be challenging for AI systems, leading to misalignment.

The alignment problem is particularly pertinent to AI models that show prominent levels of autonomy where they can learn and adapt from data without explicit instructions for every conceivable situation. Models can "misbehave" if they interpret their training data or the objective function, they are supposed to optimize in ways not intended by their human designers. The main AI types most affected by the alignment problem are shown in Table 2.

Table 2: Types of AI and their Relevance to Alignment

Type of Method	Definition	Relevance to Alignment
Deep Learning/Neural Networks:	Multi-layered artificial neural networks designed to recognize patterns in data.	Deep learning models can sometimes produce unexpected results due to training data biases or misinterpreting complex patterns. (Amodei, et al., 2016)
Reinforcement Learning	An AI method where agents learn how to respond in an environment by performing specific actions and receiving rewards or penalties	RL agents can discover "shortcuts" to achieve high rewards that do not align with the intended goals. (Hadfield-Menell, et al., 2017)
Generative Adversarial Networks (GANs):	A pair of neural networks, one generating data and the other evaluating it, trained in tandem.	GANs can produce data that maximizes the objective function but does not necessarily align with human values or expectations. (Goodfellow et al., 2014).
Evolutionary Algorithms:	Algorithms that apply the principles of natural evolution, such as mutation, crossover, and selection, to produce better solutions over generations.	These algorithms can sometimes evolve solutions that exploit loopholes or unintended features of the fitness function, leading to undesired outcomes. (Lehman, J et al., 2020).
Open-ended Learning Systems:	Systems that continue to learn and evolve without a fixed endpoint or goal.	Without concrete goals, these systems can evolve in unpredictable and potentially misaligned ways (Stanley et al., 2015).

Examples of TAP in AI-supported business processes

It is essential to mention that the alignment problem is unrestricted to these AI types. However, due to their adaptive and often "black box" nature, these methodologies have highlighted the challenge most vividly in recent AI research. Listed below are ten examples of the alignment problem in AI when used in information systems for business decision-making. These examples illustrate the many ways in which AI alignment problems can manifest in business settings. Ethical and societal impacts must be carefully considered when deploying AI systems.

1. **Bias in Recruitment Algorithms:** AI-driven recruitment tools may inadvertently perpetuate biases against certain demographic groups. For instance, an AI system trained on past hiring data might develop biases against women or minority groups (Russell et al., 2022).
2. **Credit Scoring Systems:** AI systems used in credit scoring can sometimes unfairly penalize individuals based on geographical location or past financial behavior, which may not accurately predict their creditworthiness. (Burrell, 2016).
3. **Automated Trading Algorithms:** AI in financial trading can lead to unforeseen market volatility or crashes, as these systems might make decisions that are profitable in the short term but harmful in the long term. (Lewis, 2014).
4. **Customer Service Chatbots:** AI-driven chatbots might misinterpret customer emotions or concerns, leading to customer dissatisfaction or miscommunication. (Kaplan & Haenlein, 2019).
5. **Supply Chain Optimization:** AI systems used for supply chain optimization might overemphasize cost-cutting, leading to unethical labor practices or environmental harm. (Choi et al., 2018).
6. **Predictive Policing Algorithms:** AI used in predictive policing can reinforce existing biases and lead to disproportionate targeting of specific communities. (Lum & Isaac, 2016).
7. **Healthcare Decision Support Systems:** AI systems in healthcare might make decisions based on incomplete data, leading to misdiagnosis or inappropriate treatment recommendations. (Topol 2019).
8. **Personalized Marketing Algorithms:** AI-driven personalized marketing can invade privacy or manipulate consumer behavior in unethical ways. (Zuboff, 2019).
9. **Employee Monitoring Tools:** AI systems used for employee monitoring can invade privacy and create a distrustful workplace environment. (Moore 2018).
10. **Algorithmic Decision-Making in Loan Approval:** AI systems used for loan approvals might discriminate against certain groups or individuals, leading to unfair financial exclusion. (O'Neil 2016).

Managerial decision making

Decision-making is the role of the manager. These decisions play a crucial role in supporting a business's strategic goals and the business processes that realize those strategic goals. There are distinct types of decisions within organizations. The nature and complexity of decisions vary. As a result, distinct methods of decision-making are needed. What follows are the primary types of decisions made within a business. Table 3 presents the decision type, its definition, and then an example. The following is an overview of the distinct types of Decision Types used in a managerial role. This review sets the stage for a review of the AI Methods used to support Decision Types. And eventually presenting the TAP Issues associated with these AI-supported Decision Types.

Table 3: Decision type

Decision Type	Definition	Example
Operational Decisions	These are routine decisions that occur often and are typically structured. They need a straightforward procedure or rules to guide the decision-making process. (Laudon & Laudon, 2016).	Processing an order, updating inventory records, and payroll systems.
Managerial Decisions	These are decisions made by mid-level managers that are semi-structured. They often require some human judgment, intuition, and evaluation. (O'Brien & Marakas, 2011)	Budget forecasts, sales analysis, and production scheduling.
Tactical Decisions	These decisions sit between operational and strategic decisions and are usually short-term, focused on specific departments or projects. (Power 2007)	Resource allocation for a particular project, deciding marketing strategies for a specific campaign.
Strategic Decisions	These are high-level decisions made by top managers and executives. They are unstructured and have long-term implications. (Turban et al., 2010)	Mergers and acquisitions, entering new markets, and long-term strategic goals.
Ad-hoc Analysis	These are unplanned or unexpected analyses often resulting from a specific query or problem. (Watson & Frolick1993)	Analyzing unexpected sales patterns and looking into an unusual customer complaint trend.
Collaborative Decisions	Decisions made collectively by a group of individuals, using the diverse expertise and perspectives of the group members. (Nunamaker et al., 1997)	A team using a collaborative software platform to decide on marketing strategies.
Knowledge-driven Decisions	Decisions that rely heavily on expert knowledge and deep understanding, often in specialized or technical fields. (Turban et al. 2005)	A medical diagnosis made using an expert system that incorporates medical research and patient data.
Data-Driven Decisions	Decisions are based on data analysis and interpretation, often using statistical tools, algorithms, and machine learning techniques. (Provost & Fawcett 2013)	A retailer uses machine learning to analyze customer data for personalized marketing campaigns.

AI methods can be used to implement these various decision types. The following subsection is an overview of those implementations. AI technologies have increasingly been incorporated into various decision-making processes to enhance accuracy, efficiency, and insights. Table 4 presents a breakdown of some AI types suited for the decision categories used in business decision-making.

Table 4: Types of Decisions and AI methods

Type of Decision	Type of Method	Alignment Issue	Example
<i>Operational Decisions</i>	Data Mining – (ML)	If data mining tools are not well-aligned, they might miss patterns or produce misleading insights. (Amodei et al., 2016).	In a bank's information systems process, if a data mining tool incorrectly identifies spending patterns, it might mistakenly classify legitimate transactions as fraudulent.
	Anomaly Detection	False positives or negatives can occur if models are not well-aligned with the expected anomalies. (Chandola et al., 2009).	An information system process handling e-commerce transaction might flag genuine high-value purchases as anomalies, causing customers inconvenience.
<i>Managerial Decisions</i>	Data Visualization	Misrepresentation or overemphasis of certain data aspects. (Kosara & Mackinlay 2013)	An information system process might show sales figures in a graph where the Y-axis does not start from zero, making growth seem more significant than it is.
<i>Tactical Decisions</i>	Predictive Analytics – (ML)	Misaligned models can give inaccurate predictions, leading to suboptimal decisions. (Gunning 2017)	An information system process for stock investment might predict a stock's rise based on historical data, but if misaligned, it could miss recent market trends, leading to financial losses.
<i>Strategic Decisions</i>	Predictive Modelling – (ML)	Offering over-confident or overly generalized predictions to executives. (Dietterich 2017).	An information system process might predict a steady growth in a sector based on past trends, not accounting for recent disruptions, causing executives to make misinformed strategic decisions.
	Natural Language Processing	Misinterpretation of user queries or providing information out of context. (Russell & Norvig 2022).	A manager asking an NLP-based information system process, "What were our weakest markets last year?" might get a list of markets without context or reasoning, leading to poor strategic decisions.

Type of Decision	Type of Method	Alignment Issue	Example
<i>Ad-hoc Analysis</i>	Natural Language Processing	Misunderstanding or misinterpreting user queries can lead to incorrect analyses. (Amodei & Clark 2016).	A user might query an analysis tool: "Which regions had below-average sales?" An alignment issue could lead to the tool interpreting "below-average" in a way that doesn't match the user's intent, e.g., excluding certain regions or using an incorrect time frame.
<i>Collaborative Decisions</i>	Recommendation Systems	Systems might make biased recommendations based on historical data or could perpetuate existing biases in group decisions. (Datta et al., 2015).	A collaborative tool might consistently suggest male candidates for a job role in tech because, historically, more men have been hired in those roles.
	Sentiment Analysis	The system might misinterpret sentiments, especially in diverse groups with different linguistic nuances. (Kiritchenko & Mohammad 2018).	In an international collaboration tool, the system might misinterpret a phrase considered positive in one culture but neutral or negative in another.
<i>Data-driven Decisions:</i>	Deep Learning	These models, particularly neural networks, are often seen as "black boxes" and can be hard to interpret. Misalignment can arise when we can't fully understand or explain decisions made by the model. (Lipton 2016)	A deep learning model might classify specific medical images as indicative of a disease, but doctors might need clarification on why, making it hard to trust the model.

Conclusions

In conclusion, this document has explored the profound implications of integrating Artificial Intelligence (AI) into managerial decision-making processes, focusing on the alignment problem. AI, defined as the simulation of human intelligence by machines, encompasses a broad spectrum of methods, including machine learning algorithms, natural language processing, recommendation systems, sentiment analysis, data visualization, anomaly detection, expert systems, neural networks, and deep learning. These technologies offer remarkable capabilities in enhancing business operations, from predictive analytics and customer service to supply chain optimization and personalized marketing.

However, the alignment problem presents significant challenges. Ensuring that AI systems act in ways that align with human values, ethics, and intentions is a multifaceted issue involving the definition of goals, value misalignment, robustness and safety, long-term consequences, and the interpretation of human preferences. This problem is particularly pertinent to AI models with elevated levels of autonomy, such as deep learning, reinforcement learning, generative adversarial networks, evolutionary algorithms, and open-ended learning systems. The potential for these systems to produce unintended behaviours underscores the need for not just design but careful design and oversight.

The practical examples offered illustrate the varied ways in which alignment problems can manifest in business settings. From biases in recruitment algorithms and credit scoring systems to challenges in automated trading, customer service chatbots, and predictive policing, it is not just clear, but crucial that ethical and societal impacts must be carefully considered when deploying AI systems.

Finally, examining managerial decision-making highlighted the distinct types of decisions—operational, managerial, tactical, strategic, ad-hoc analysis, collaborative, knowledge-driven, and data-driven—and how AI methods support these decision types. The associated alignment issues emphasize the need for robust, transparent, and interpretable AI systems to ensure that decisions made with AI support are ethical, fair, and aligned with organizational goals. While AI holds tremendous potential for transforming business processes and decision-making, addressing the alignment problem is crucial for realizing its benefits responsibly and ethically. Future research and development should continue to focus on creating AI systems that enhance efficiency and accuracy and align with human society's complex and nuanced values.

Future Research

What is next for this project is to find traditional IS systems – Transaction Processing Systems, Decision Support Systems, et al. that currently support the Business Decision Making process and further determine to what degree AI methods are used to support these systems. Understanding AI techniques' role in supporting these systems will help find alignment problems. Another extension of this project would be figuring out the effect of replacing traditional IS systems with pure AI technology. For example, strategic decision support software is being replaced by an LLM like ChatGPT, which uses precise prompts to create strategic decisions. This investigation would be followed by developing guidelines for avoiding or reducing alignment issues when using AI-supported Information Systems processes to make business decisions.

References

- Amodei, D., Olah, C., Steinhardt, J., Christiano, P., Schulman, J., & Mané, D. (2016). "Concrete Problems in AI Safety." arXiv preprint arXiv:1606.06565
- Bostrom, N. (2014). *Superintelligence: Paths, Dangers, Strategies*. Oxford University Press.
- Burrell, J. (2016). How the machine 'thinks': Understanding opacity in machine learning algorithms. *Big Data & Society*, 3(1), 2053951715622512
- Chandola, V., Banerjee, A., & Kumar, V. (2009). Anomaly detection: A survey. *ACM computing surveys (CSUR)*, 41(3), 1-58
- Chen, M., Mao, S., & Liu, Y. (2014). Big data: A survey. *Mobile Networks and Applications*, 19(2), 171-209
- Choi, T. M., Wallace, S. W., & Wang, Y. (2018). Big data analytics in operations management. *Production and Operations Management*, 27(10), 1868-1883
- Christian, B. (2020). *The Alignment Problem: Machine Learning and Human Values*. W. W. Norton & Company.
- Datta, A., Tschantz, M. C., & Datta, A. (2015). Automated experiments on ad privacy settings. *Proceedings on Privacy Enhancing Technologies*, 2015(1), 92-112
- Dietterich, T. G. (2017). Steps toward robust artificial intelligence. *AI Magazine*, 38(3), 3-24
- Duan, Franklin L. "Artificial Intelligence." 2023, https://doi.org/10.1007/978-981-19-8394-8_2.

- Goodfellow, I., Pouget-Abadie, J., Mirza, M., Xu, B., Warde-Farley, D., Ozair, S., ... & Bengio, Y. (2014). Generative adversarial nets. In *Advances in Neural Information Processing Systems* (pp. 2672-2680).
- Gunning, D. (2017). Explainable artificial intelligence (XAI). Defense Advanced Research Projects Agency (DARPA), Web
- Hadfield-Menell, D., Milli, S., Abbeel, P., Russell, S. J., & Dragan, A. (2017). Inverse reward design. In *Advances in Neural Information Processing Systems* (pp. 6765-6774).
- Kaplan, A., & Haenlein, M. (2019). Siri, Siri, in my hand: Who's the fairest in the land? On the interpretations, illustrations, and implications of artificial intelligence. *Business Horizons*, 62(1), 15-25.
- Kiritchenko, S., & Mohammad, S. M. (2018). Examining gender and race bias in two hundred sentiment analysis systems. *arXiv preprint arXiv:1805.04508*.
- Kosara, R., & Mackinlay, J. (2013). Storytelling: The next step for visualization. *Computer*, (5), 44-50.
- Laudon, K. C., & Laudon, J. P. (2016). *Management information system: managing the digital firm*. Pearson/Prentice Hall
- Lehman, J., Clune, J., Misevic, D., Adami, C., Altenberg, L., Beaulieu, J., ... & Blount, Z. (2020). The surprising creativity of digital evolution: A collection of anecdotes from evolutionary computation and artificial life research communities. *Artificial Life*, 26(2), 274-306.
- Lewis, M. (2014). *Flash Boys: A Wall Street Revolt*. W. W. Norton & Company.
- Lipton, Z. C. (2016). The mythos of model interpretability. *arXiv preprint arXiv:1606.03490*.
- Liu, B. (2012). *Sentiment Analysis and Opinion Mining*. Morgan & Claypool Publishers.
- Lum, K., & Isaac, W. (2016). To predict and serve? *Significance*, 13(5), 14-19
- Moore, P. V. (2018). *The quantified self in precarity: Work, technology, and what counts*. Routledge.
- Nunamaker, J. F., Jr., Briggs, R. O., & Mittleman, D. D. (1997). Electronic Meeting Systems: Ten Years of Lessons Learned. In D. Coleman & R. Khanna (Eds.), *Groupware: Technology and Applications* (pp. 149-193). Prentice Hall.
- O'Brien, J. A., & Marakas, G. M. (2011). *Management Information Systems*. McGraw-Hill/Irwin.
- O'Neil, C. (2016). *Weapons of math destruction: How big data increases inequality and threatens democracy*. Crown.
- Power, D. J. (2007). *A Brief History of Decision Support Systems*. DSSResources.COM
- Provost, F., & Fawcett, T. (2013). *Data Science for Business: What You Need to Know about Data Mining and Data-Analytic Thinking*. O'Reilly Media.
- Ricci, F., Rokach, L., & Shapira, B. (2011). *Introduction to Recommender Systems Handbook*. Springer US.
- Russell, S. J. (2019). *Human Compatible: Artificial Intelligence and The Problem of Control*. Viking.
- Russell, S. J., & Norvig, P. (2022). *Artificial intelligence: a modern approach*. 4th ed. Malaysia; Pearson Education Limited
- Stanley, K. O., & Lehman, J. (2015). *Why greatness cannot be planned: The myth of the objective*. Springer

- Topol, E. J. (2019). *Deep medicine: How artificial intelligence can make healthcare human again*. Basic Books
- Tufte, E. R. (2001). *The Visual Display of Quantitative Information* (2nd ed.). Graphics Press.
- Turban, E., Aronson, J. E., & Liang, T. (2005). *Decision Support Systems and Intelligent Systems* (7th ed.). Prentice Hall.
- Turban, E., Sharda, R., & Delen, D. (2010). *Decision Support and Business Intelligence Systems*. Prentice Hall
- Watson, H. J., & Frolick, M. N. (1993). Determining information requirements for an EIS. *MIS Quarterly*, 17(3), 327-343
- Zuboff, S. (2019). *The age of surveillance capitalism: The fight for a human future at the new frontier of power*. PublicAffair