

The Seven Tools of Causal Inference with Reflections on Machine Learning

JUDEA PEARL, UCLA Computer Science Department, USA

Systems that operate in purely statistical mode of inference entail theoretical limits on their power and performance. Such systems cannot reason about interventions and retrospection and, therefore, cannot serve as the basis for strong AI. To achieve human-level intelligence, learning machines need the guidance of a model of external reality, similar to the ones used in causal inference tasks. To demonstrate the essential role of such models, this paper presents a summary of seven tasks which are beyond reach of associational learning systems and which have been accomplished using the tools of causal modeling.

ACM Reference Format:

Judea Pearl. 2018. The Seven Tools of Causal Inference with Reflections on Machine Learning. 1, 1 (July 2018), 6 pages. <https://doi.org/10.1145/nnnnnnn.nnnnnnn>

1 INTRODUCTION

The dramatic success in machine learning has led to an explosion of AI applications and increasing expectations for autonomous systems that exhibit human-level intelligence. These expectations, however, have met with fundamental obstacles that cut across many application areas. One such obstacle is *adaptability* or *robustness*. Machine learning researchers have noted that current systems lack the capability of recognizing or reacting to new circumstances they have not been specifically programmed or trained for. Intensive theoretical and experimental efforts toward “transfer learning,” “domain adaptation,” and “Lifelong learning” [Chen and Liu 2016] are reflective of this obstacle.

Another obstacle is *explainability*, that is, “machine learning models remain mostly black boxes” [Ribeiro et al. 2016] unable to explain the reasons behind their predictions or recommendations, thus eroding users trust. and impeding diagnosis and repair. See [Marcus 2018] and (<http://www.sciencemag.org/news/2018/05/ai-researchers-allege-machine-learning-alchemy>).

A third obstacle concerns the understanding of cause-effect connections. This hallmark of human cognition [Lake et al. 2015; Pearl and Mackenzie 2018] is, in this author’s opinion, a necessary (though not sufficient) ingredient for achieving human-level intelligence. This ingredient should allow computer systems to choreograph a parsimonious and modular representation of their environment, interrogate that representation, distort it by acts of imagination and finally answer “What if?” kind of questions. Examples are interventional questions: “What if I make it happen?” and retrospective or explanatory questions: “What if I had acted differently?” or “what if my light had not been late?”

Author’s address: Judea Pearl, UCLA Computer Science Department, 4532 Boelter Hall, Los Angeles, California, 90095-1596, USA, judea@cs.ucla.edu.

© 2018 Association for Computing Machinery.

This is the author’s version of the work. It is posted here for your personal use. Not for redistribution. The definitive Version of Record was published in , <https://doi.org/10.1145/nnnnnnn.nnnnnnn>.

I postulate that all three obstacles mentioned above require equipping machines with causal modeling tools, in particular, causal diagrams and their associated logic. Advances in graphical and structural models have made counterfactuals computationally manageable and thus rendered causal reasoning a viable component in support of strong AI.

In the next section, I will describe a three-level hierarchy that restricts and governs inferences in causal reasoning. The final section summarizes how traditional impediments are circumvented using modern tools of causal inference.

THE THREE LAYER CAUSAL HIERARCHY

A useful insight unveiled by the theory of causal models is the classification of causal information in terms of the kind of questions that each class is capable of answering. The classification forms a 3-level hierarchy in the sense that questions at level i ($i = 1, 2, 3$) can only be answered if information from level j ($j \geq i$) is available.

Figure 1 shows the 3-level hierarchy, together with the characteristic questions that can be answered at each level. The levels are titled 1. Association, 2. Intervention, and 3. Counterfactual. The names of these layers were chosen to emphasize their usage. We call the first level Association, because it invokes purely statistical relationships, defined by the naked data.¹ For instance, observing a customer who buys toothpaste makes it more likely that he/she buys floss; such association can be inferred directly from the observed data using conditional expectation. Questions at this layer, because they require no causal information, are placed at the bottom level on the hierarchy. The second level, Intervention, ranks higher than Association because it involves not just seeing what is, but changing what we see. A typical question at this level would be: What will happen if we double the price? Such questions cannot be answered from sales data alone, because they involve a change in customers behavior, in reaction to the new pricing. These choices may differ substantially from those taken in previous price-raising situations. (Unless we replicate precisely the market conditions that existed when the price reached double its current value.) Finally, the top level is called Counterfactuals, a term that goes back to the philosophers David Hume and John Stewart Mill, and which has been given computer-friendly semantics in the past two decades. A typical question in the counterfactual category is “What if I had acted differently,” thus necessitating retrospective reasoning.

Counterfactuals are placed at the top of the hierarchy because they subsume interventional and associational questions. If we have a model that can answer counterfactual queries, we can also answer questions about interventions and observations. For example, the interventional question, What will happen if we double the

¹Some other terms used in connection to this layer are: “model-free,” “model-blind,” “black-box,” or “data-centric.” Darwiche [2017] used “function-fitting,” for it amounts to fitting data by a complex function defined by the neural network architecture.

Level (Symbol)	Typical Activity	Typical Questions	Examples
1. Association $P(y x)$	Seeing	What is? How would seeing X change my belief in Y ?	What does a symptom tell me about a disease? What does a survey tell us about the election results?
2. Intervention $P(y do(x), z)$	Doing Intervening	What if? What if I do X ?	What if I take aspirin, will my headache be cured? What if we ban cigarettes?
3. Counterfactuals $P(y_x x', y')$	Imagining, Retrospection	Why? Was it X that caused Y ? What if I had acted differently?	Was it the aspirin that stopped my headache? Would Kennedy be alive had Os- wald not shot him? What if I had not been smoking the past 2 years?

Fig. 1. The Causal Hierarchy. Questions at level i can only be answered if information from level i or higher is available.

price? can be answered by asking the counterfactual question: What would happen had the price been twice its current value? Likewise, associational questions can be answered once we can answer interventional questions; we simply ignore the action part and let observations take over. The translation does not work in the opposite direction. Interventional questions cannot be answered from purely observational information (i.e., from statistical data alone). No counterfactual question involving retrospection can be answered from purely interventional information, such as that acquired from controlled experiments; we cannot re-run an experiment on subjects who were treated with a drug and see how they behave had they not been given the drug. The hierarchy is therefore directional, with the top level being the most powerful one.

Counterfactuals are the building blocks of scientific thinking as well as legal and moral reasoning. In civil court, for example, the defendant is considered to be the culprit of an injury if, *but for* the defendant's action, it is more likely than not that the injury would not have occurred. The computational meaning of *but for* calls for comparing the real world to an alternative world in which the defendant action did not take place.

Each layer in the hierarchy has a syntactic signature that characterizes the sentences admitted into that layer. For example, the association layer is characterized by conditional probability sentences, e.g., $P(y|x) = p$ stating that: the probability of event $Y = y$ given that we observed event $X = x$ is equal to p . In large systems, such evidential sentences can be computed efficiently using Bayesian Networks, or any number of machine learning techniques.

At the interventional layer we find sentences of the type $P(y|do(x), z)$, which denotes "The probability of event $Y = y$ given that we intervene and set the value of X to x and subsequently observe event $Z = z$. Such expressions can be estimated experimentally from randomized trials or analytically using Causal Bayesian Networks [Pearl 2000, Chapter 3]. A child learns the effects of interventions through playful manipulation of the environment (usually in a deterministic playground), and AI planners obtain interventional knowledge by

exercising their designated sets of actions. Interventional expressions cannot be inferred from passive observations alone, regardless of how big the data.

Finally, at the counterfactual level, we have expressions of the type $P(y_x|x', y')$ which stand for "The probability that event $Y = y$ would be observed had X been x , given that we actually observed X to be x' and Y to be y' . For example, the probability that Joe's salary would be y had he finished college, given that his actual salary is y' and that he had only two years of college." Such sentences can be computed only when we possess functional or Structural Equation models, or properties of such models [Pearl 2000, Chapter 7].

This hierarchy, and the formal restrictions it entails, explains why machine learning systems, based only on associations, are prevented from reasoning about actions, experiments and causal explanations.²

THE SEVEN TOOLS OF CAUSAL INFERENCE (OR WHAT YOU CAN DO WITH A CAUSAL MODEL THAT YOU COULD NOT DO WITHOUT?)

Consider the following five questions:

- How effective is a given treatment in preventing a disease?
- Was it the new tax break that caused our sales to go up?
- What is the annual health-care costs attributed to obesity?
- Can hiring records prove an employer guilty of sex discrimination?
- I am about to quit my job, but should I?

The common feature of these questions is that they are concerned with cause-and-effect relationships. We can recognize them through words such as "preventing," "cause," "attributed to," "discrimination," and "should I." Such words are common in everyday language, and

²Some readers have pointed out to me that the limitation of level-1 of the hierarchy can perhaps be alleviated by current deep learning techniques, which try to minimize "over fit," in contrast to classical statistical techniques which try to maximize "fit." Unfortunately, the theoretical barriers that separate the three layers in the hierarchy tell us that the nature of our objective function does not matter. As long as our system optimizes some property of the observed data, however noble or sophisticated, while making no reference to the world outside the data, we are back to level-1 of the hierarchy with all the limitations that this level entails.

our society constantly demands answers to such questions. Yet, until very recently science gave us no means even to articulate them, let alone answer them. Unlike the rules of geometry, mechanics, optics or probabilities, the rules of cause and effect have been denied the benefits of mathematical analysis.

To appreciate the extent of this denial, readers would be stunned to know that only a few decades ago scientists were unable to write down a mathematical equation for the obvious fact that “mud does not cause rain.” Even today, only the top echelon of the scientific community can write such an equation and formally distinguish “mud causes rain” from “rain causes mud.”

Things have changed dramatically in the past three decades, A mathematical language has been developed for managing causes and effects, accompanied by a set of tools that turn causal analysis into a mathematical game, not unlike solving algebraic equations, or finding proofs in high-school geometry. These tools permit us to express causal questions formally, codify our existing knowledge in both diagrammatic and algebraic forms, and then leverage our data to estimate the answers. Moreover, the theory warns us when the state of existing knowledge or the available data are insufficient to answer our questions; and then suggests additional sources of knowledge or data to make the questions answerable.

This development has had a transformative impact on all data-intensive sciences, especially social science and epidemiology, in which causal diagrams have become a second language [Morgan and Winship 2015; VanderWeele 2015]. In these disciplines, causal diagrams have helped scientists extract causal relations from associations, and de-construct paradoxes that have baffled researchers for decades [Pearl and Mackenzie 2018; Porta 2014].

I call the mathematical framework that led to this transformation “Structural Causal Models (SCM).”

The SCM deploys three parts

- (1) Graphical models,
- (2) Structural equations, and
- (3) Counterfactual and interventional logic

Graphical models serve as a language for representing what we know about the world, counterfactuals help us to articulate what we want to know, while structural equations serve to tie the two together in a solid semantics.

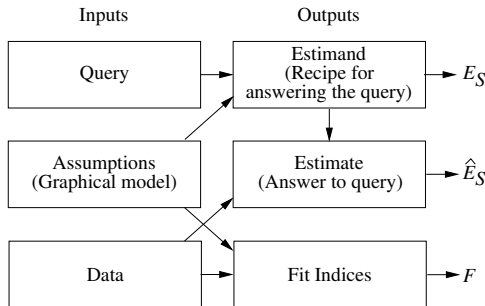
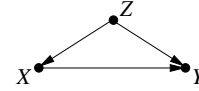


Fig. 2. How the SCM “inference engine” combines data with causal model (or assumptions) to produce answers to queries of interest.

Figure 2 illustrates the operation of SCM in the form of an inference engine. The engine accepts three inputs: Assumptions, Queries, and Data, and produces three outputs: Estimand, Estimate and Fit indices. The Estimand (E_S) is a mathematical formula that, based on the Assumptions, provides a recipe for answering the Query from any hypothetical data, whenever they are available. After receiving the Data, the engine uses the Estimand to produce an actual Estimate ($\hat{E}_S$) for the answer, along with statistical estimates of the confidence in that answer (To reflect the limited size of the data set, as well as possible measurement errors or missing data.) Finally, the engine produces a list of “fit indices” which measure how compatible the data are with the Assumptions conveyed by the model.

To exemplify these operations, let us assume that our Query stands for the causal effect of X (taking a drug) on Y (recovery), written $Q = P(Y|do(X))$. Let the modeling assumptions be encoded in the graph below where Z is a third variable (say Gender) affecting



both X and Y . Finally, let the data be sampled at random from a joint distribution $P(X, Y, Z)$. The Estimand (E_S) calculated by the engine (using Tool 2, below) will be the formula $E_S = \sum_z P(Y|X, Z)P(Z)$, which defines a procedure of estimation. It calls of estimating the gender-specific conditional distributions $P(Y|X, Z)$ for males and females, weighing them by the probability $P(Z)$ of membership in each gender, then taking the average. Note that the estimand E_S defines a property of $P(X, Y, Z)$ that, if properly estimated, would provide a correct answer to our Query. The answer itself, the Estimate $\hat{E}_S$, can be produced by any number of techniques that produce a consistent estimate of E_S from finite samples of $P(X, Y, Z)$. For example, the sample average (of Y) over all cases satisfying the specified X and Z conditions, would be a consistent estimate. But more efficient estimation techniques can be devised to overcome data sparsity [Rosenbaum and Rubin 1983]. This is where deep learning techniques excel, and where they are often employed [van der Laan and Rose 2011].

Finally, the Fit Index in our example will be NULL. In other words, after examining the structure of the graph, the engine should conclude (using Tool 1, below) that the assumptions encoded do not have any testable implications. Therefore, the veracity of the resultant estimate must lean entirely on the assumptions encoded in the graph – no refutation nor corroboration can be obtained from the data.³

The same procedure applies to more sophisticated queries, for example, the counterfactual query $Q = P(y_x|x', y')$ discussed before. We may also permit some of the data to arrive from controlled experiments, which would take the form $P(V|do(W))$, in case W is the controlled variable. The role of the Estimand would remain that of converting the Query into the syntactic format of the available

³The assumptions encoded in the graph are conveyed by its missing arrows. For example, Y does not influence X or Z , X does not influence Z and, most importantly, Z is the only variable affecting both X and Y . That these assumptions lack testable implications can be concluded from the fact that the graph is complete, i.e., no edges are missing.

data and, then, guiding the choice of the estimation technique to ensure unbiased estimates. Needless to state, the conversion task is not always feasible, in which case the Query will be declared “non-identifiable” and the engine should exit with FAILURE. Fortunately, efficient and complete algorithms have been developed to decide identifiability and to produce estimands for a variety of counterfactual queries and a variety of data types [Bareinboim and Pearl 2016; Shpitser and Pearl 2008; Tian and Pearl 2002].

Next we provide a bird’s eye view of seven tasks accomplished by the SCM framework, the tools used in each task, and discuss the unique contribution that each tool brings to the art of automated reasoning.

Tool 1: Encoding Causal Assumptions – Transparency and Testability

The task of encoding assumptions in a compact and usable form, is not a trivial matter once we take seriously the requirement of transparency and testability.⁴ Transparency enables analysts to discern whether the assumptions encoded are plausible (on scientific grounds), or whether additional assumptions are warranted. Testability permits us (be it an analyst or a machine) to determine whether the assumptions encoded are compatible with the available data and, if not, identify those that need repair.

Advances in graphical models have made compact encoding feasible. Their transparency stems naturally from the fact that all assumptions are encoded graphically, mirroring the way researchers perceive of cause-effect relationship in the domain; judgments of counterfactual or statistical dependencies are not required, since these can be read off the structure of the graph. Testability is facilitated through a graphical criterion called *d*-separation, which provides the fundamental connection between causes and probabilities. It tells us, for any given pattern of paths in the model, what pattern of dependencies we should expect to find in the data [Pearl 1988].

Tool 2: Do-calculus and the control of confounding

Confounding, or the presence of unobserved causes of two or more variables, has long been considered the the major obstacle to drawing causal inference from data. This obstacle had been demystified and “deconfounded” through a graphical criterion called “back-door.” In particular, the task of selecting an appropriate set of covariates to control for confounding has been reduced to a simple “roadblocks” puzzle manageable by a simple algorithm [Pearl 1993].

For models where the “back-door” criterion does not hold, a symbolic engine is available, called *do-calculus*, which predicts the effect of policy interventions whenever feasible, and exits with failure whenever predictions cannot be ascertained with the specified assumptions [Pearl 1995; Shpitser and Pearl 2008; Tian and Pearl 2002].

Tool 3: The Algorithmization of Counterfactuals

Counterfactual analysis deals with behavior of specific individuals, identified by a distinct set of characteristics. For example, given that

Joe’s salary is $Y = y$, and that he went $X = x$ years to college, what would Joe’s salary be had he had one more year of education.

One of the crowning achievements of modern work on causality has been to formalize counterfactual reasoning within the graphical representation, the very representation researchers use to encode scientific knowledge. Every structural equation model determines the truth value of every counterfactual sentence. Therefore, we can determine analytically if the probability of the sentence is estimable from experimental or observational studies, or combination thereof [Balke and Pearl 1994; Pearl 2000, Chapter 7].

Of special interest in causal discourse are counterfactual questions concerning “causes of effects,” as opposed to “effects of causes.” For example, how likely it is that Joe’s swimming exercise was a necessary (or sufficient) cause of Joe’s death [Halpern and Pearl 2005; Pearl 2015a].

Tool 4: Mediation Analysis and the Assessment of Direct and Indirect Effects

Mediation analysis concerns the mechanisms that transmit changes from a cause to its effects. The identification of such intermediate mechanism is essential for generating explanations and counterfactual analysis must be invoked to facilitate this identification. The graphical representation of counterfactuals enables us to define direct and indirect effects and to decide when these effects are estimable from data, or experiments [Pearl 2001; Robins and Greenland 1992; VanderWeele 2015]. Typical queries answerable by this analysis are: What fraction of the effect of X on Y is mediated by variable Z .

Tool 5: Adaptability, External Validity and Sample Selection Bias

The validity of every experimental study is challenged by disparities between the experimental and implementational setups. A machine trained in one environment cannot be expected to perform well when environmental conditions change, unless the changes are localized and identified. This problem, and its various manifestations are well recognized by AI researchers, and enterprises such as “domain adaptation,” “transfer learning,” “life-long learning,” and “explainable AI” [Chen and Liu 2016], are just some of the subtasks identified by researchers and funding agencies in an attempt to alleviate the general problem of robustness. Unfortunately, the problem of robustness, in its broadest form, requires a causal model of the environment, and cannot be handled at the level of Association. Associations are not sufficient for identifying the mechanisms affected by changes that occurred [Pearl and Bareinboim 2014]. The reason being that surface changes in observed associations do not uniquely identify the underlying mechanism responsible for the change. The *do-calculus* discussed above now offers a complete methodology for overcoming bias due to environmental changes. It can be used both for re-adjusting learned policies to circumvent environmental changes and for controlling bias due to non-representative samples [Bareinboim and Pearl 2016].

⁴Economists, for example, having chosen algebraic over graphical representations, are deprived of elementary testability-detecting features [Pearl 2015b].

Tool 6: Recovering from Missing Data

Problems of missing data plague every branch of experimental science. Respondents do not answer every item on a questionnaire, sensors malfunction as weather conditions worsen, and patients often drop from a clinical study for unknown reasons. The rich literature on this problem is wedded to a model-free paradigm of associational analysis and, accordingly, it is severely limited to situations where missingness occurs at random, that is, independent of values taken by other variables in the model. Using causal models of the missingness process we can now formalize the conditions under which causal and probabilistic relationships can be recovered from incomplete data and, whenever the conditions are satisfied, produce a consistent estimate of the desired relationship [Mohan and Pearl 2018; Mohan et al. 2013].

Tool 7: Causal Discovery

The d -separation criterion described above enables us to detect and enumerate the testable implications of a given causal model. This opens the possibility of inferring, with mild assumptions, the set of models that are compatible with the data, and to represent this set compactly. Systematic searches have been developed which, in certain circumstances, can prune the set of compatible models significantly to the point where causal queries can be estimated directly from that set [Pearl 2000; Peters et al. 2017; Spirtes et al. 2000].

Alternatively, Shimizu et al. [2006] proposed a method of discovering causal directionality based on functional decomposition [Peters et al. 2017]. The idea is that in a linear model $X \rightarrow Y$ with non-Gaussian noise, $P(y)$ is a convolution of two non-Gaussian distributions and would be, figuratively speaking, “more Gaussian” than $P(x)$. The relation of “more Gaussian than” can be given precise numerical measure and used to infer directionality of certain arrows.

Tian and Pearl [2002] developed yet another method of causal discovery based on the detection of “shocks,” or spontaneous local changes in the environment which act like “Nature’s interventions,” and unveil causal directionality toward the consequences of those shocks.

CONCLUSIONS

I have argued that causal reasoning is an indispensable component of human thought that should be formalized and algorithmized toward achieving human-level machine intelligence. I have explicated some of the impediments toward that goal in the form of a three-level hierarchy, and showed that inference to levels 2 and 3 require a causal model of one’s environment. I have described seven cognitive tasks that require tools from those two levels of inference and demonstrated how they can be accomplished in the SCM framework.

It is important to note that the models used in accomplishing these tasks are structural (or conceptual), and requires no commitment to a particular form of the distributions involved. On the other hand, the validity of all inferences depend critically on the veracity of the assumed structure. If the true structure differs from the one

assumed, and the data fits both equally well, substantial errors may result, which can only be assessed through a sensitivity analysis.

It is also important to keep in mind that the theoretical limitations of model-free machine learning do not apply to tasks of prediction, diagnosis and recognition, where interventions and counterfactuals assume a secondary role.

However, the model-assisted methods by which these limitations are circumvented can nevertheless be applicable to problems of opacity, robustness, explainability and missing data, which are generic to machine learning tasks. Moreover, given the transformative impact that causal modeling has had on the social and medical sciences, it is only natural to expect a similar transformation to sweep through the machine learning technology, once it is enriched with the guidance of a model of the data-generating process. I expect this symbiosis to yield systems that communicate with users in their native language of cause and effect and, leveraging this capability, to become the dominant paradigm of next generation AI.

ACKNOWLEDGMENTS

This research was supported in parts by grants from Defense Advanced Research Projects Agency [#W911NF-16-057], National Science Foundation [#IIS-1302448, #IIS-1527490, and #IIS-1704932], and Office of Naval Research [#N00014-17-S-B001]. This paper benefited substantially from comments by anonymous reviewers and conversations with Adnan Darwiche.

REFERENCES

- A. Balke and J. Pearl. 1994. Probabilistic evaluation of counterfactual queries. In *Proceedings of the Twelfth National Conference on Artificial Intelligence*. Vol. I. MIT Press, Menlo Park, CA, 230–237.
- E. Bareinboim and J. Pearl. 2016. Causal inference and the data-fusion problem. *Proceedings of the National Academy of Sciences* 113 (2016), 7345–7352. Issue 27.
- Z. Chen and B. Liu. 2016. *Lifelong Machine Learning*. Morgan and Claypool Publishers, San Rafael, CA.
- A. Darwiche. 2017. *Human-Level Intelligence or Animal-Like Abilities?* Technical Report. Department of Computer Science, University of California, Los Angeles, CA. arXiv:1707.04327.
- J.Y. Halpern and J. Pearl. 2005. Causes and Explanations: A Structural-Model Approach—Part I: Causes. *British Journal of Philosophy of Science* 56 (2005), 843–887.
- B.M. Lake, R. Salakhutdinov, and J.B. Tenenbaum. 2015. Human-level concept learning through probabilistic program induction. *Science* 350 (2015), 1332–1338. Issue 6266.
- G. Marcus. 2018. *Deep Learning: A Critical Appraisal*. Technical Report. Departments of Psychology and Neural Science, New York University, NY. (<https://arxiv.org/pdf/1801.00631.pdf>).
- K. Mohan and J. Pearl. 2018. *Graphical Models for Processing Missing Data*. Technical Report R-473, (http://ftp.cs.ucla.edu/pub/stat_ser/r473.pdf). Department of Computer Science, University of California, Los Angeles, CA. Forthcoming, *Journal of American Statistical Association* (JASA).
- K. Mohan, J. Pearl, and J. Tian. 2013. Graphical Models for Inference with Missing Data. In *Advances in Neural Information Processing Systems 26*, C.J.C. Burges, L. Bottou, M. Welling, Z. Ghahramani, and K.Q. Weinberger (Eds.). 1277–1285. (<http://papers.nips.cc/paper/4899-graphical-models-for-inference-with-missing-data.pdf>).
- S.L. Morgan and C. Winship. 2015. *Counterfactuals and Causal Inference: Methods and Principles for Social Research (Analytical Methods for Social Research)* (2nd ed.). Cambridge University Press, New York, NY.
- J. Pearl. 1988. *Probabilistic Reasoning in Intelligent Systems*. Morgan Kaufmann, San Mateo, CA.
- J. Pearl. 1993. Comment: Graphical Models, Causality, and Intervention. *Statist. Sci.* 8, 3 (1993), 266–269.
- J. Pearl. 1995. Causal diagrams for empirical research. *Biometrika* 82, 4 (1995), 669–710.
- J. Pearl. 2000. *Causality: Models, Reasoning, and Inference*. Cambridge University Press, New York. 2nd edition, 2009.
- J. Pearl. 2001. Direct and indirect effects. In *Uncertainty in Artificial Intelligence, Proceedings of the Seventeenth Conference*. Morgan Kaufmann, San Francisco, CA, 411–420.

- J. Pearl. 2015a. Causes of Effects and Effects of Causes. *Journal of Sociological Methods and Research* 44 (2015), 149–164. Issue 1.
- J. Pearl. 2015b. Trygve Haavelmo and the emergence of causal calculus. *Econometric Theory* 31 (2015), 152–179. Issue 1. Special issue on Haavelmo Centennial.
- J. Pearl and E. Bareinboim. 2014. External validity: From *do*-calculus to transportability across populations. *Statist. Sci.* 29 (2014), 579–595. Issue 4.
- J. Pearl and D. Mackenzie. 2018. *The Book of Why: The New Science of Cause and Effect*. Basic Books, New York.
- J. Peters, D. Janzing, and B. Schölkopf. 2017. *Elements of Causal Inference – Foundations and Learning Algorithms*. The MIT Press, Cambridge, MA.
- M. Porta. 2014. The deconstruction of paradoxes in epidemiology. (2014). OUPblog, <https://blog.oup.com/2014/10/deconstruction-paradoxes-sociology-epidemiology/> (May 14).
- M.T. Ribeiro, S. Singh, and C. Guestrin. 2016. Why should I trust you?: Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. ACM, New York, NY, 1135–1144.
- J.M. Robins and S. Greenland. 1992. Identifiability and Exchangeability for Direct and Indirect Effects. *Epidemiology* 3, 2 (1992), 143–155.
- P. Rosenbaum and D. Rubin. 1983. The central role of propensity score in observational studies for causal effects. *Biometrika* 70 (1983), 41–55.
- S. Shimizu, P.O. Hoyer, A. Hyvärinen, and A.J. Kerminen. 2006. A linear non-Gaussian acyclic model for causal discovery. *Journal of the Machine Learning Research* 7 (2006), 2003–2030.
- I. Shpitser and J. Pearl. 2008. Complete Identification Methods for the Causal Hierarchy. *Journal of Machine Learning Research* 9 (2008), 1941–1979.
- P. Spirtes, C.N. Glymour, and R. Scheines. 2000. *Causation, Prediction, and Search* (2nd ed.). MIT Press, Cambridge, MA.
- J. Tian and J. Pearl. 2002. A general identification condition for causal effects. In *Proceedings of the Eighteenth National Conference on Artificial Intelligence*. AAAI Press/The MIT Press, Menlo Park, CA, 567–573.
- Mark J. van der Laan and Sherri Rose. 2011. *Targeted Learning: Causal Inference for Observational and Experimental Data*. Springer, New York.
- T.J. VanderWeele. 2015. *Explanation in Causal Inference: Methods for Mediation and Interaction*. Oxford University Press, New York.