

The Eight Pillars of Causal Wisdom

by Judea Pearl

An edited transcript of a lecture given at the
West Coast Experiments, April 2017

video posted at <https://www.youtube.com/watch?v=8nHVUFqI0zk>

Thank you for getting up so early in the morning, it is barely midnight in my calendar. I appreciate your coming to to this talk which I prepared with great difficulty. Until today I have been talking and lecturing about various approaches, nuances and jargons within the causal inference community. Today I want to talk about what this community can say to the rest of the scientific world. I want to represent all of you and ask: What have we contributed to science since day one.

What is “day one?”

I take “day one” to be Sewall Wright’s paper of 1920’s in which he drew the first causal diagram, which he called a “path diagram.”

I would like to explain why I consider that paper to be “day one.” I am a computer scientist, and computer scientists are very sensitive to notation and to languages. Primarily because we have to teach this stuff to a stupid robot and, if you do not have the right notation you do not convey the right information to the robot and things can get really scary. It is for this reason that my criterion for separating one idea from another revolves around notation, not the idea itself (which may be subject to interpretation) – how you express the idea, what notation you use, and what grammar regulates your expressions. This is also my perception of human thought. Human language is a faithful representative of, and in fact a guardian and controller of human thought.

Okay, with that philosophical introduction, let’s get to what I planned to say.

I want to talk about eight pillars of causal wisdom and the reasons I chose this title are:

1. To counter Steven Stigler (The slide misspells his name – it should Stigler, not Sigler) who wrote a popular and influential book on the history of Statistics, titled “The Seven Pillars of Statistics Wisdom.” I want to challenge Stigler’s title with the thesis that if you do not include causality then statistics is one pillar short of wisdom. In Stigler’s book we find only two references to causality and none to modern causal analysis.
2. Next, I want to challenge a recent paper by Angrist and Pischke (2017), in which they propose a model-blind approach to econometric education, a proposal which I believe to be misguided if not un-doable.
3. Thirdly, I want to present some ideas for computer scientists who are enthusiastic about deep learning. Research in causal inference has brought to the surface solid theoretical

limitations on what one can do and cannot do with statistically based methods of machine learning and big data. Most researchers in machine learning are not aware of those limitations.

4. Finally, and this is my main goal, I would like to empower students, faculty and all members of this creative community with a sense of accomplishment, and awareness of what this community has achieved since Sewall Wright 1920's paper. It is sort of a progress report of a century of work in – in what? In a discipline we call “causal inference.”

What makes causal inference a separate discipline? After all, every field of science is using some cause and effect relationships. Now here I come and I pick up bits and pieces from various sciences and call it “a field” – is that legitimate? Yes, it's a “field,” if you look at it from a computer science viewpoint. Those bits and pieces are united through a common and distinct grammar that is not used elsewhere. That is what makes it absolutely unique.

I'll now go over the eight pillars of wisdom which are listed right here. I'll go quickly through them – just reading – because I'm going to come back and discuss them one by one.

Pillar 1: Graphical models for prediction and diagnosis – (The list, by the way, is not chronological and I'm not going to name all the people who contributed to the results listed, I wish to focus on the results, which should make each one in this community very very proud) So prepare yourself to be proud.

Pillar 2: The control of confounding – a done problem.

Pillar 3: *Do*-calculus – An all-seeing oracle for predicting the effects of policies and interventions. I know that the skeptics among you may say: no, we haven't gotten to the end of the policy-making problem, I will be happy to discuss it in the QA session.

Pillar 4: The Algorithmization of counterfactuals,

Pillar 5: Mediation analysis and the assessment of direct and indirect effect.

Pillar 6: External validity – half a century of threats which have turned recently into opportunities

Pillar 7: Missing data, and

Pillar 8: Causal discovery.

I would like you to consider a list of five typical causal problems - we frequently encounter such problems in our research and in everyday lives. They share one feature in common, they are all using causal grammar.

Here they are:

1. How effective is a given treatment in preventing a disease.

2. Was it the new tax break that caused our sales to go up or our marketing campaign?
3. What is the annual health care cost attributed to obesity?
4. Can hiring records prove an employer guilty of sex discrimination? And the last one is on personal decision making
5. I just quit my job. Will I regret it?

As you see, I marked in yellow the words that characterize each of the five questions as “causal,” because you couldn’t articulate them even two or three decades ago. They were not subject to mathematical analysis. Fisher couldn’t even utter them (mathematically speaking). They were not part of standard science. Oddly, science has been very un-generous to causal reasoning. It totally deprived causal thinking from the benefits of mathematical notation, mathematical machinery and mathematical tools. Science has developed beautiful mathematics for optics, geometry, even probability – but, not for cause-effect relationship. That neglect has been corrected lately, and we should all be both grateful and proud of the correction. We can now write a formula for each one of these problems. And once you have a formula, you have the freedom to use all the machineries that mathematics offers us, to combine these questions with data to see how they relate to each other, and what internal constraints the grammar imposes. So, it’s a gift. It’s a gift that was given to science or to humanity, if you want to be really spiritual about it.

I will now try to explain why my teachers could not provide a formula for the sentence “mud does not cause rain.” And when I say “my teachers” I really mean our teachers. Yes, until very very recently, our teachers and professors could not write down a formula for the simple fact that “mud does not cause rain” or that “the rooster crow does not cause the sun rise” or that “the falling barometer does not cause the incoming storm.” Such sentences were considered beyond mathematics, and here are a few reasons why this happened.

First of all, we have a linguistic mismatch. Algebraic equations are symmetrical and here we are talking about asymmetric relationship. A second reason – causal thinking was meticulously purged from statistics, by Galton and Pearson – deliberately and explicitly (See for instance Pearson’s Grammar of Science, 1911.) The statistical education that emerged became an education without causality (I sometimes call it causal-a-phobia.) And the grammar produced by Sewall Wright, which was a diagrammatic abstraction of structural equations was totally misunderstood, unappreciated, and under-developed until the 1980s. Arthur Goldberger called it “A scandal of neglect.” And it got misinterpreted so badly that, today, even in economics, structural equations are an embarrassment to education. (See Chris Auld paper, this conference.)

The interpretation, by the way, is very simple, though I don’t find it in any econometric textbook. The interpretation of structural equations is simply “a society of listening agents” I say it here to all of you who are supposed to know it, and probably know it, in the same spirit that Ed Leamer said it in 1985 – he noticed that every econometrician knows what endogenous variables are and what structural equations are etc. etc., but he couldn’t find a single econometric textbook that defines it properly.

Here is the definition in one sentence: “A society of listening variables,” – each variable listens and responds to others. It is very important to use the word “listen” because it’s asymmetric. If I listen to you, it doesn’t mean that you listen to me.

Everything follows from that metaphor – everything. Compare

$$y = \beta x + u \text{ vs. } y \longleftarrow \beta x + u$$

The difference between the left-hand side (ie, a regression equation with an algebraic equality sign) and the right-hand side which is a structural equation is that the left one represents DATA and the right one represents REALITY. And that's why I feel upset when I see universities like Columbia building a whole building for "data science" and not "reality science." You probably have a department in your university dedicated to data science. But what about reality? What I really mean is that "data science" should include the interpretation of data, not merely the summarization of data.

Okay, if we start talking history, why not go ten thousand years ago? Something happened ten thousand years ago in which causal inference played a major role. Ten thousand years ago human beings accounted for less than one-tenth of 1 percent of all vertebrate life on planet earth. Today that percentage including livestock and pets is in the neighborhood of 98 percent. (This is taken from Daniel Dennett's book) and the question is, what happened? When computer scientists ask "what happened," what they mean is "what computational facilities did human acquire 10,000 years that they did not possess before?" And the answer is partly here: At about 70,000 years ago sapiens from East Africa moved into the Arabian Peninsula and from there they quickly overran the entire Eurasian landmass, wiping out the native population there – Homo Erectus in Asia and the Neanderthal man in Europe. What caused that transition? It is conjectured that the secret of success was the ability of Homo-Sapiens to do counterfactual reasoning, i.e., to imagine the impossible. This ivory figurine is the first artifact found which proves people's ability to imagine things that do not exist in the universe. Not a copy of prior observations, but a totally new object.

Here it is magnified. It is called "the Lion Man of Stadel Cave," found in a cave in Germany, and is dated 32,000 years ago. It is the first combination of things that do not combine in our natural environment. There isn't such a creature that is half man and half lion, it was nevertheless put together by an artist to convey a new idea. If you take two things that do not exist together and you put them together, the combination may evoke a new idea: perhaps that I want to be as bold as a lion or as strong as a lion. This imaginative ability, according to Harari's book Sapiens, accounted for the development of larger organizations and larger communities and all the powers and benefits that come with larger communities. So, for instance, one tribe could have promised another tribe protection on the basis of a unrealizable yet communicated promise that protection will prevail. It then generated a "market for promises." "I promise you something that you'll get tomorrow and you've got to trust me that I will deliver on my promise." That market of promises enabled a chief to unite members of the tribe and motivate them to do things that they wouldn't do otherwise. And that's unique to human beings. For instance (quoting Harari), you could never convince a monkey to give you a banana by promising him limitless bananas after death in monkey heaven. Oddly, people do go for such stream of bananas.

When you go through the trouble of formal analysis, and ask: what is special about counterfactuals? you come to the realization that we have here a hierarchy of cognitive functions, or a "ladder of causation." Here is the ladder. On the lowest rung you can

see the Neanderthal man, owls and snakes – These creatures developed through a slow evolutionary process resembling machine learning today. Snakes and owls can spot a prey from a few miles away, and perform miracles that we cannot duplicate in any laboratory today. But they cannot invent eyeglasses or telescopes. Going up the ladder we find the ability to predict the result of DOING, and that corresponds to the second rung of the ladder which is called “intervention.” The top level, where you can see Einstein and the invention of airplanes and other beautiful technologies involves imagination of putting together things that never existed before. It involves imagining combinations, and asking hypothetical questions such as “would Kennedy be alive if Oswald had not shot him?” Or, What if I had not smoked for the last two years? It is a quest for explanation, retrospective thinking, which sits at the top of the ladder.

Analysis shows that you cannot get answers to questions at level i unless you have information from level i or higher. One simple manifestation of this constraint is the mantra of “correlation does not imply causation.” Indeed, correlation resides on level one – association – and if you’re asking questions about causation or policy intervention, you must have some information from level two or higher.

I will demonstrate this ladder in a problem that economists should be familiar with – supply and demand – which I’m going to display right here. It is taken from Goldberger’s model of supply and demand. You have the price, P , the demand quantity Q , and you have income I , and wages W . Income and Price affect the quantity demanded, while Wages and Demand affect the price. So you have here equilibrium between the manufacturer who asked himself “what price should I set for my product tomorrow?,” and the consumer who is asking: “how much can I buy given the price?”

I posed this question to over a hundred economists. That was in 1998, when I was trying to peddle my book *Causality* and I thought economists would be very interested in this exercise. So I asked these three questions about the model. All economists were perfectly proficient in estimating the coefficients b_1, b_2, d_1 and d_2 from data. They are all identifiable. Ed says no, I think they are, because W and I act as instrumental variables. But this is all standard econometrics, and people can do all kind of tricks to identify the coefficients. My question was not about the coefficients, my question was about what you want to do with them, once they are estimated. So I asked, for instance, “What is the expected value of the demand Q if currently the price is reported to be p_0 ?” The next question has to do with intervention – “What is the expected value of the demand if the price is set (by government or some other agency or by whimsical decision maker) at p_0 ?” And the third question was about counterfactuals: “Given that the current price is p_0 , what would the expected value of the demand be if we were to set the price at p_1 ?” Going back in time.

Note that I’m not talking about solving here, only about expressing what needs to be estimated. Well, it turned out that even economists found it very hard to express the question mathematically, except the first one. The first question was purely statistical, hence very easy (bing) “the expected value of Q given that you observed the price P .” Everybody answered it. I then asked the second question and no one touched it, except for one professor, Ed Leamer, who said it obviously has to do with the coefficient b_1 , which is defined as the causal effect of P on Q .

As to the third question, no one could write down an expression for the quantity needed. Again, I am not talking about solving it, just writing down an expression, which

requires counterfactual notation. You currently observe the price at p_0 and you're asking retrospectively, what would Q be had the price actually been p_1 instead of p_0 . Solving it is another question. It turns out all three questions are solvable within the linear system. They have a simple and elegant solution – yet no one could handle it, at least in 1998. Today, I hope that every economist can do it half asleep. [Yes, I know that I'm over optimistic but it is better than being a fatalistic.]

Okay, now that we come to this point, we can talk about the language needed to express causal problems. Notice that I've already used two languages. Unwittingly, I slipped in on you two languages; one is the language of the model – structural equations and diagram. The diagram, of course, is just a decorative way of expressing some abstraction of the structural equation model. [You agree! I got one agreement here!] Okay, so why two languages? One is a diagram or structural equations, and other one is the language in which we express what we wish to be estimated. Two different species. And look, they look so foreign to each other. The first has arrows or equations, okay? And the second, look at it, expected values of counterfactuals, subscripts, “do-operators” - crazy.

Why am I insisting of a bilingual approach? This is what I'm going to explain on the next slide. Yes, I'm advocating a bilingual approach to causal reasoning. We need to have a special language to explicitly specify what we know and separate it from the language we use to specify what we wish to know. I call the second quantity “a query.” A query is a different species from what we know, ie, from the knowledge. I'll give you an example. Well, this is the flow chart of the logic. The red signifies what the user or the researcher needs to specify, which includes three things, without which you cannot proceed. You've got to specify your problem, so this is your query [right here]. Next, you have to tell me what is to be estimated. You would think it's obvious, but not in the papers that I need to review. The research problem is the most neglected part in the papers that I happened to review, and you have probably experienced the same thing. Authors can write pages upon pages on what he or she has done, while research problems are left for the reviewers to uncover.

Next, you need specify what knowledge you have. In the next slide, I'll tell you why we need to have a separate language for that. Why we cannot use the same language to specify both the query and the knowledge. Finally, you need to have some data. Note, the researcher is not exonerated from the task of specifying what data he wants to include in the analysis. Do you want to include experimental studies or only observational studies or, maybe run experiments on another variable that you can manipulate, an IV, for instance. So you have to specify what kinds of data are available to you. Some people call it: “design,” but all it is is specifying in some language what sort of data are available.

Good. Now I'm going to add to it one more channel called “testability,” which is required if you want to test your model against the data. For model testing, you should feed the graph and the data into a box called “testability” and assess the degree of fit. They do it routinely in structural equation models, but it is rarely done in economics (as far as I was able to tell). The ability to identify the testable implications of a causal model is one of the virtues of graphical models.

The items in green exemplify the three specifications: Data, Knowledge and Query. For instance, your data may be from an observational study, that is, the probability of x , y , and z . Your state of knowledge may be described by this diagram: X causes Z . Z causes Y ,

maybe you have some unmeasured variable a confounder U , okay. Note that you specify knowledge in the way that is stored in your mind. He who asks you to labor and torture yourself before you can express what you know is doing you disservice and diminishes the veracity of your judgment. In most cases, knowledge is stored qualitatively in your mind, in terms of cause-effect relationships, put it down on paper!!, it is the most reliable way of specifying knowledge. If you want to be quantitative, you'll have a chance. Now examine the query, it is an expression involving a *do*-operator, if you're trying to analyze interventions, or counterfactual notation, if you want to do retrospective thinking.

Now, who is going to put them all together? What is the relationship between this crazy graph here and those fancy expressions, which involve counterfactuals or interventions? Who makes sure that the two cohere with each other. This is a job of the structural equation models and this is what I am putting on top, here. It's not a new invention. It is written into the fundamental idea of structural equations except for few details. Unfortunately, these details are not recognized in your field, so, you should be the one to spread the word. The key word is that every structural equation model dictates counterfactuals. Every structural equation model assigns a probability to every conceivable counterfactual sentence, no matter how complex or convoluted. Once we accept that fact, it becomes clear that we're not talking about two foreign objects – graphs and counterfactuals – we're talking about two aspects of the object – structural equation model.

Now I'm going to justify why we need a bilingual approach and why we cannot do it in the most powerful one of the two. namely, the counterfactual approach. Here is a simple example, I presented it last year in this conference, but I don't think it sunk through, because I am still hearing colleagues telling me: "I'm a potential outcome person" – as if they are exonerated from any of these specifications. Or, as if it is possible to do everything in the potential outcome language. Okay, let's go and try to do it.

Here is a simple toy story, involving smoking, tar and cancer. Smoking causes tar accumulation in your lungs and this affects whether or not you're going to get cancer within three years. Plus, you might have some unobserved genetic trait, or genotypes that drives you to nicotine and at the same time put you in a cancer risk. It's a very simple story. I just used 45 seconds to convey it to you and you understood it? It was conveyed in English, and this is the amazing power of natural languages. Now we want to formalize it and I'm going to present to you two representation of the same story; the are equivalent, one is written in the jargon of potential outcome and the second is written in the jargon of graphs. And we're going to judge them not just by their aesthetic, but by concrete criteria. Here is the representation of the problem in counterfactual language. Take my word, I've gone through it several times, each one of these equations is necessary and sufficient.

This is how it looks if you really want to do it with counterfactual and forget about graphs and arrows. You cannot do it with fewer equations. The guy who promised me to think about it was Paul Holland. I posed this in a seminar in Berkeley in 1993. He promised me an answer and he presented it in 1995 at a conference. I'm still waiting for the answer. They simply cannot do that – even at the level of representation – look at this representation. This is how the world looks in their eyes – including the eyes of many of us who are guilty of saying "I'm a potential outcome person," without internalizing what it entails.

I'm going now to ask some questions about this model. If this is the world in the eyes of

the potential outcome camp, it would be interesting to know what one do with it and what one can't.

1. Consistency – Can you look at those assumptions and tell me if they are consistent? i.e., perhaps one of them violates the others. It's doable, believe me. There is a mathematical logic that can handle counterfactuals. It's elaborate, but it can be done formally, and the answer is: Yes, these statements are consistent! None of them violates the others. But can you verify it directly from the assumptions as they are written here? Hard, isn't it? In fact it is not doable by any mortal that I know.
2. Are they complete? Namely, perhaps we've forgotten to write down one assumption which is part of the story. I gave you the story, and you understood it. Can you tell whether I have forgotten any assumption that is necessary for replicating the story? As you can see, it's s hard to tell from this syntax.
3. Are they redundant? Perhaps one of the assumptions follows from the others. Hard to tell. As I mentioned, there is a logic for checking redundancy, but it is not obvious.
4. Are they plausible? Which means, do they match the story. The story is plausible, right? But if you just look at that representation, can you tell whether all these assumptions are compatible with your experience with, or your knowledge about cancer and smoking. Again, it's hard to tell.
5. Are they testable? This one is quite concrete. Is there any statistical method by which we can test them? Or, are there any statistical data that could possibly violate those assumption?

All these questions are very hard to answer in this representation. Now look at the alternative. in the form of structural equations. You simply draw the story the way I told it to you, specifying "Who affects whom" "Where you have a confounder?" [It is right there]. etc etc, This graph is your final specification which is exactly equal to the counterfactual specification above it. Exactly. If you want me, I can write those counterfactual sentences straight from the graph. There is a back-and-forth mapping between the two; you give me those counterfactuals, and I can draw the graph.

Most remarkably, once I have the graph I can answer questions about consistency, plausibility, redundancy and testability by inspection. The answers are transparent.

This difference in transparency is not merely a matter of convenience. but reflects a fundamental impediment. In computer science, we know that the choice of representation is the key to complexity. Tasks that take polynomial time in one representation may take exponential time in another. Even though there is a one-to-one correspondence between the two. This is what I mean when I call structural equation "an oracle"; you do not need to explicitly specify all those counterfactuals when you build a model, they are implicit, and available upon demand.

I finished the introduction. And now we go back to the Eight Pillars of Causal Wisdom
After working in this area for maybe twenty years, one owes society a progress report. So here is our progress report: Eight Pillars. This progress is the work of an entire community. I didn't put names on the slide, because I want to focus on the results and their

significance. This is what has been produced with just one simple idea; that the syntax of cause effect relationships deserves a special attention. It's a different kind of syntax. It's a different kind of thinking, hence it requires a different language. And, behind the language, it deserves semantics and, behind that, it deserves mathematical machinery to derive things which did not exist before Sewall Wright put down an arrow on paper.

How much time to I have? Good, one minute per pillar.

[Pillar] 1. In the 1980's there was a development called Bayesian networks in computer science, at the time when we were looking for expert systems to replicate doctors and lawyers by machines. And it turned out that the graphical model called Bayesian network performed quite well on this task. You have a network there which enables you to update beliefs on the basis of evidence quickly, swiftly, and reliably and using message passing scheme. I consider that to be a gift of the Gods that probability theory found a way to represent itself in graphs. It's a gift from God because whoever expected the axioms of conditional independence to share such a large core with the axioms of graph interception. And they do! Interception in graphs is governed by 4-5 axioms that happen to coincide with Dawid's axioms of conditional independence. That was a gift from God because it enabled us to use "*d*-separation" – the trace that reality (or its causal model) leaves in our data. It enables you to look at the graph and decide, without thinking of Dawid's axioms, which variables are conditional independent given others.

Pillar 2 also deserves a minute of reflection. The menace of confounding is now de-confounded. I know it sounds like a sweeping statement, but I am prepared to defend it. This menace has been around since Pearson discovered "spurious correlations" (1896) and took it as a proof that correlation is not causation and, eventually, that causation is superfluous. It was rediscovered by Yule and by Simpson, It has been a topic of contention in epidemiology from the day that epidemiologists looked at data. I'm saying that confounding is totally de-confounded. And my proof is that we can decide, looking at the graph, what covariates need to be controlled for, if we want to eliminate confounding. The same applies if you do matching, or if you want to do propensity score – all these methods are just different names for procedures that people have devised to fight confounding, but one thing is key to all these procedures and one puzzle that is avoided in the literature that deals with these methods: what covariates should we adjust for. Well, the answer is given to you by the "backdoor criterion," which is a graphical equivalent of what the potential outcome people call "strong ignorability."

This criterion answers many questions that economists ask themselves on a daily basis. For example, which coefficient in an economic model is identified by OLS? Imagine a structural equation model containing 47 variables and a mishmash of arrows going all over the place. We are asking: Which of the model's coefficients is estimable by OLS? This question has a very simple answer, requiring only that we look at the graph and examine the arrows going around each of the coefficients.

Many econometric questions can be answered by the same graphical criterion, including questions of endogeneity, identification, testability and robustness. I know that none of you teaches this method, and that is why I am mentioning it – I want you and your students to know what you're missing.

Okay, Pillar 3 takes us beyond adjustment. It asks the more general question: When can the effect of intervention be estimated by any analytical method, without physically

intervening. Here we are aiming beyond adjustment. And the answer lies in a game called “*do*-calculus” which tells you whether or not your policy expression, which involves *do*-operators, can be reduced to a *do*-free expression. It is a game of logic. If you apply the rules correctly, you will change the expression and if you succeed in eliminating the *do*-operator, what do you have? You have only “see-operators” and what is “see?” “See” means that we found an estimand based on observational studies. That’s exactly what we mean by identification. Done. Complete. Rest. And this can be done in polynomial time.

Pillar 4: The algorithmization of counterfactuals. I touched on that already. Counterfactuals are not the product of a whimsical mind nor relationships between treatment and outcome but a feature of physical reality as represented by structural equations. This is what I meant by saying that every structural equation model assigns a probability measure to every counterfactual sentence. Therefore, if people share a model of the world, they should agree on all counterfactuals.

This explains why we agree on the statement that “if Oswald did not kill Kennedy, then somebody else did,” and we don’t agree on the sentence “if Oswald were not to kill Kennedy somebody else would have.” We agree on the first and not the second one and we form a consensus. How can you explain that we form a consensus on something as fanciful as 1963 sentiments in Dallas Texas? Well, that explains it, because we share a perception of what causes what in the world and that dictates our agreement on counterfactuals. Examples of policy-relevant counterfactuals are: (1) ETT – the effect of treatment on the treated, and (2) Necessary and sufficient causes of a given effect. These are counterfactual tasks that are constantly being analyzed in the literature, and we now have a mathematical handle on their estimability

[Pillar] 5. The next one: Mediation analysis – A nice pillar. Everybody wants to understand the mechanism by which a cause transmits changes to its effects, and counterfactual logic must be invoked if you want to find things like indirect effects. The graphical representation enables us to decide when direct and indirect effects are estimable from observational or experimental data or various combinations of the two. This is also analyzable in polynomial time.

Pillar 6. External validity and selection bias. Elias is going to elaborate on this topic on Tuesday. I believe Deaton is also going to talk about it today. I just like to say that what we are seeking here are guarantees that “what works here would also work elsewhere” (Deaton and Cartwright). It is a very reasonable and simple requirement. However, since Campbell and Stanley coined the word “external validity” in 1963, the field has merely accumulated threats to that requirement, threats upon threats and the list is increasing with the imagination of the people who try to solve the problem. We haven’t seen any mechanism for disarming those threats. With the work of Elias, we have today a solid theory that tells you when you can disarm threat 1 or threat 2, whether you have enough knowledge to do that, and what do you need to do to ensure that “what works here would also work elsewhere.” It’s all formulated mathematically.

So I’m very happy with Pillar 6 and I’m going to Pillar 7 – Missing data – which Karthika is going to talk about on Tuesday. That’s another area which we used to think lies totally within the province of statistics. Not so! It is actually a causal problem. Now that we understand the interplay between different grammars, we can tell whether a task belongs in one province or another. The idea is that you cannot recover missing data without

assumptions about the reasons for the missingness. “Reasons for missingness” means causal modeling. And if you model the problem causally and use all the machinery that the causal-inference community has accumulated, you can answer beautiful questions about missing data. For example, can you recover the missing data? Do you have a consistent estimate of a needed parameter? Note that the question of whether a consistent estimate exists or doesn’t exist has not been asked before. Even in the statistical community because they could not articulate the assumptions that are needed for a solution. It is a peculiar yet recurrent phenomenon in science. You don’t ask questions that you cannot solve.

Okay, the last pillar is “Causal Discovery,” which I’m not going to elaborate on. I hope Clark Glymour will tell us more about it. It deals with ways of recovering, or discovering the causal graph behind the data. It is based on mild yet universal assumptions of “faithfulness” or “stability.” I’m running out of time.

Before I finish, I would like however to quote and re-quote the statement by Gary King which says:

“More has been learned about causal inference in the last few decades than the sum total of everything that has been learned about it in all prior recorded history.”

It is a very sweeping statement. That is why I quote it again and again. And I will now finish with Democritos who invited me to be a king of Persia – a job that I’m not sure my wife would like me to take, given the political climate. I would still like to discover one causal relation.

Thank you very much for being so attentive. Credit goes to the people and teams listed here and to many more whom I did not mention. Many people participated and contributed ideas to these developments. References can be found in my survey papers, for example, http://ftp.cs.ucla.edu/pub/stat_ser/r416-reprint.pdf.

Thank you very much.

I’m open to questions of course. Catch me day and night.