

Intervention with Limited Information^{1,2}

Drew Fudenberg³ and David K. Levine⁴

Abstract

We study how optimal interventions in response to a shock with limited information depend on the complexity of the system. We show that as the complexity of the system grows, the optimal intervention shrinks to zero.

¹We are grateful to Chris Ackerman, Daniel Clark, Daniel Friedman, George Georgiadis, Giacomo Lanzani, Andy McLennan, John Rust and Larry Samuelson for helpful comments, and to NSF grant SES 1643517 and MIUR (PRIN 2017H5KPLL) for financial support.

²This version: 10/1/2021. First version: 08/22/2020

³Department of Economics, MIT; drew.fudenberg@gmail.com

⁴Department of Economics and RSCAS, EUI, and Department of Economics, WUSTL; david@dklevine.com

1. Introduction

We model the problem of a decision maker faced with limited information who has the opportunity to intervene in a system that has received a shock as the choice of a location in n -dimensional Euclidean space. The loss is measured by the Euclidean distance from the optimum, which was at 0 before the shock. The decision maker knows how large the shock is, but not the direction in which the optimum has moved. They do, however, have partial information about this: they know a direction in which the loss is decreasing. In this setting complexity is measured by the dimension of the system n . We show that in a one-dimensional system the full optimum is obtained, and that as the dimension increases optimal interventions become smaller and the optimized loss grows. The intuition is that with more directions it is more likely that an intervention in a randomly chosen direction will “overshoot” and lead to an additional loss. In the limit as $n \rightarrow \infty$ it is optimal not to intervene at all and simply accept the loss. In addition, following Stokey (2009), if there is a small fixed cost of making adjustments, then it is best not to respond to the shock at all. Moreover, in more complex systems the shock must be bigger before it is worthwhile to intervene.

Imagine, for example, that while driving a car you hear a large bang, the car no longer accelerates properly, and the engine starts to make loud noises. Should you get out of the car and try to fix it, or simply struggle on? You are not an automobile mechanic, and do not know how engines work, but you are aware of basic facts, for example, that breaking parts of the engine with a hammer will make things worse. In settings of aiding others there is a substantial literature arguing that non-intervention is generally thought to be better.⁵ However, we expect that someone’s perceived obligation to intervene depends on the amount of information they have relative to the difficulty of the problem. For example, in a health emergency on an airplane a doctor would be expected to intervene when an ordinary passenger would not. By contrast if an elderly person drops a bag of groceries even an ordinary passer-by might be expected to help. Similar issues arise with respect to economic systems: How much should policy makers intervene in response to a pandemic? How should rich countries intervene to help a developing country facing a crisis?⁶ How should business firms faced with unexpected systems failures or by competition from new products respond?

Our result also has implications for the evolution of a firm, organization, or organism: as the environment grows more complex, evolution will slow to a halt unless the available information becomes richer. This is because our result shows that with limited information a single “experiment” or “mutation” is likely to yield little gain over the status quo.⁷

⁵Spranca, Minsk and Baron (1991) and Cushman, Young and Hauser (2006), for example, provide experimental evidence that this is a widely held view and review other evidence and literature.

⁶Easterly (2002) discusses some of the ways such interventions can backfire.

⁷We are grateful to a referee for this interpretation.

The problem of intervention with limited information as we formulate it appears to be new. It is connected to Samuelson (1947)’s LeChatelier principle, which shows that in a setting of certainty, where the location of the optimum is known, the size of intervention is strictly smaller when intervention is possible in a subset of the dimensions instead of all of them. We extend this to the case where the location of the optimum is unknown, and show that this effect increases as the complexity of the system does. Our setup has also some similarities to Ely (2011), who also uses the number of control dimensions as a measure complexity. Our results are also, in a certain sense, a counterpoint to the “curse of dimensionality” (see, for example, Bellman (1957)) that states that the time it takes to reach the optimum increases exponentially with the dimension of the problem. Here we show that incomplete optimization also does less well as the dimension of the problem increases. We relate this to results about the rate of convergence of gradient descent algorithms.

2. The Model

A decision maker must choose an action x in $\mathbb{R}^n$ to minimize a quadratic loss function $(1/2) \sum_{i=1}^n (x_i - \hat{x}_i)^2$. The decision maker knows the magnitude $|\hat{x}|$ of the bliss point $\hat{x}$, but not its location: They know only that it is drawn from a distribution that is spherically symmetric around 0, and that there is a direction d in which the loss is decreasing. It is best to choose x lying in direction d , so the problem of the decision maker reduces to choosing a non-negative real number ϕ and setting $x = \phi d$. The corresponding loss is $(1/2)|\hat{x} - \phi d|^2$, so the expected loss is this amount integrated with respect to $\hat{x}$ according to a uniform distribution over the surface of the n -dimensional half-sphere $x_1 \geq 0$ of radius $|\hat{x}|$.

One way to think about one-dimensional adjustment is to imagine the system as a mechanical machine controlled by dials, with each dial corresponding to a coordinate axis. Being restricted to a single dimension corresponds to partial ignorance, in the sense of being aware of only one dial. A decision maker who knew the system better might be aware of more dials/dimensions. For example, in trying to help a person with a severe leg wound a typical person might think that putting pressure on the wound is a good idea, while a doctor might know also that it is possible to use a tourniquet above the wound (a second dial).

The assumption that intervention is only contemplated in a single dimension is a good approximation in many settings with partial ignorance. For example, the EU debate over pandemic aid focused on a single dimension: the division of aid between grants and loans. A less well known example is that of the Wright brothers. An airplane is a complex system, and marketing one is complicated as well. Faced with competition from Glenn Curtiss’ airplanes built using the superior aileron technology, the Wrights did not consider redesigning their airplane or change their marketing practices: the only dimension in which they

responded was in the filing of patent lawsuits.⁸

It should be clear as well that some tinkering in other dimensions does take place. For example, in the EU pandemic package, the pro-grant countries agreed to decrease the total size of the aid package and increase the rebates of the “Frugal Four.” These other dimensions typically are limited to small changes. Subsequently we consider the robustness of our results to the possibility that the decision maker might have some idea of a slightly better direction than d .

Preliminary Analysis

The problem is homogeneous, so the optimal solution $\hat{\phi}^n$ is independent of $|\hat{x}|$, and if the minimum expected loss corresponding to $|\hat{x}| = 1$ is L^n , the corresponding expected loss for general $|\hat{x}|$ is $L^n|\hat{x}|^2$. With this in mind, we normalize $|\hat{x}| = 1$. It is convenient moreover to use coordinates in which $\hat{x} = (1, 0, 0, \dots, 0)^T$, so that d rather than $\hat{x}$ is random and uniformly distributed over the unit half-sphere $d_1 \geq 0$. We subsequently study this normalized problem.

In the one dimensional case, $n = 1$, the true optimum lies at $x_1 = 1$ so that the optimal choice of ϕ is $\hat{\phi}^1 = 1$ with corresponding loss $L^1 = 0$. The key point is that in higher dimensions there is a tradeoff. If $d = \hat{x}$ it would be best to choose $\phi = 1$. If d is orthogonal to $\hat{x}$ it would be best to choose $\phi = 0$. Hence the optimal choice of ϕ is a compromise: a large ϕ works well for “good” directions r close to $\hat{x}$ but poorly in “bad” directions far from $\hat{x}$. The intuition we would like to establish is that in higher dimensions there are relatively more “bad” directions so smaller ϕ should be chosen.

3. The Main Result

Theorem 1. *The optimal solution $\hat{\phi}^n$ and corresponding loss L^n satisfies $\hat{\phi}^{n+1} < \hat{\phi}^n$, $L^{n+1} > L^n$ with $\hat{\phi}^1 = 1$, $L^1 = 0$ and $\lim_{n \rightarrow \infty} \hat{\phi}^n = 0$, $\lim_{n \rightarrow \infty} L^n = 1/2$.*

This says that given a shock that without intervention would result in a half unit loss, if the system has high complexity, as measured by n , it is better to intervene very little with the result that almost the entire loss must be swallowed. If the system has low complexity, a substantial intervention is optimal, resulting in a substantial mitigation of the loss.

The case $n = 1$ was proven above. The remainder of the proof follows from several intermediate results which we now prove.

Lemma 2. *The loss function is $(1/2)(1 - 2a^n\phi + (\phi)^2)$ with $0 < a^n \leq 1$ and $a^1 = 1$.*

Proof. Fix d . Using the fact that $|d| = 1$ we can compute the loss as $L^n(d, \phi) = (1/2)[(1 - \phi d_1)^2 + \sum_{i=2}^n (\phi d_i)^2] = (1/2)[1 - 2d_1\phi + \sum_{i=1}^n (\phi d_i)^2] = 1/2 - d_1\phi + (\phi)^2/2$. Hence a^n is the integral of d_1 with respect to d which is distributed uniformly over the unit half-sphere $|d| = 1$ and $d_1 \geq 0$. As d_1 is strictly positive

⁸A good account of the sad saga of the Wright brothers can be found in Shulman (2002).

with probability 1 and $d_1 \leq 1$ on the unit half-sphere it follows that $0 < a^n \leq 1$, and the result for $n = 1$ is immediate. $\square$

Corollary 3. $\phi^n = a^n$ and $L^n = 1/2 - (a^n)^2/2$.

The theorem will follow if we can show that $a^{n+1} < a^n$ and $\lim_{n \rightarrow \infty} a^n = 0$. This we show next.

Lemma 4. Let $h(n, \theta)$ be the density function

$$h(n, \theta) = \frac{(\sin \theta)^{n-2}}{\int_0^{\pi/2} (\sin \omega)^{n-2} d\omega}$$

on $[0, \pi/2]$. Then for $n > 1$ we have $a^n = \int_0^{\pi/2} h(n, \theta) \cos \theta d\theta$.

Proof. For $n > 1$ we do the integration by taking $r_1 = \cos \theta$ where $\theta \in [0, \pi/2]$. For given θ the density is given by the surface area S^{n-2} of the $n-2$ dimensional sphere with radius equal to r_1 , which is $S^{n-2} = c(n-2)(\sin \theta)^{n-2}$ where $c(n-2)$ is positive, known, and irrelevant. Hence we can compute

$$a^n = \frac{\int_0^{\pi/2} c(n-2)(\sin \theta)^{n-2} \cos \theta d\theta}{\int_0^{\pi/2} c(n-2)(\sin \theta)^{n-2} d\theta}.$$

$\square$

The final step of proving the theorem is then

Lemma 5. $a^{n+1} < a^n$ and $\lim_{n \rightarrow \infty} a^n = 0$.

Proof. To show $a^{n+1} < a^n$ note that since $\cos \theta$ is strictly decreasing⁹ in θ it suffices to prove that $h(n+1, \theta)$ first order stochastically dominates $h(n, \theta)$. Observe that if $\theta' > \theta$ since $\sin \theta$ is strictly increasing we have

$$\frac{h(n+1, \theta')}{h(n+1, \theta)} = \left(\frac{\sin \theta'}{\sin \theta} \right)^{n-1} > \left(\frac{\sin \theta'}{\sin \theta} \right)^{n-2} = \frac{h(n, \theta')}{h(n, \theta)}.$$

Hence as both are density functions it must be that $h(n+1, 0) < h(n, 0)$ and there is a unique value $\theta \in [0, \pi/2]$ where $h(n+1, \theta) = h(n, \theta)$ which implies stochastic dominance.

To show $\lim_{n \rightarrow \infty} a^n = 0$, let $\varepsilon > 0$, and observe that there is $\theta_\varepsilon \in [0, \pi/2]$ such that for all $\theta \in [\theta_\varepsilon, \pi/2]$, $\cos(\theta) < \varepsilon$. Moreover, since $\sin \theta$ is strictly increasing on $[0, \pi/2]$ for any $M \in (\theta_\varepsilon, \pi/2)$ we have

$$\lim_{n \rightarrow \infty} \frac{\int_0^{\theta_\varepsilon} c(n-2)(\sin \theta)^{n-2} d\theta}{\int_0^{\pi/2} c(n-2)(\sin \theta)^{n-2} d\theta} \leq \lim_{n \rightarrow \infty} \frac{\theta_\varepsilon (\sin \theta_\varepsilon)^{n-2}}{\int_M^{\pi/2} (\sin \theta)^{n-2} d\theta} \leq \lim_{n \rightarrow \infty} \frac{\theta_\varepsilon (\sin \theta_\varepsilon)^{n-2}}{(\pi/2 - M)(\sin M)^{n-2}} = 0$$

⁹Note that the only facts we use about $\cos \theta$ is that on $[0, \pi/2]$ it is strictly decreasing from one to zero.

Since $\cos \theta \leq 1$

$$\lim_{n \rightarrow \infty} a_n \leq \lim_{n \rightarrow \infty} \frac{\int_0^{\theta_\epsilon} c(n-2)(\sin \theta)^{n-2} d\theta}{\int_0^{\pi/2} c(n-2)(\sin \theta)^{n-2} d\theta} + \frac{\int_{\theta_\epsilon}^{\pi/2} c(n-2)(\sin \theta)^{n-2} \epsilon d\theta}{\int_0^{\pi/2} c(n-2)(\sin \theta)^{n-2} d\theta} \leq \epsilon.$$

Since ϵ was arbitrary it follows that in fact $\lim_{n \rightarrow \infty} a^n = 0$. $\square$

Remark. We include a direct proof that $\lim_{n \rightarrow \infty} a^n = 0$ since it is straightforward, but it follows from more general known results. The fact that $\lim_{n \rightarrow \infty} a^n = 0$ follows from the fact that the first coordinate of a uniform distribution on the sphere converges in probability to 0, which in turn follows from a stronger result attributed to Poincaré¹⁰: The first coordinate of a uniform distribution on an n dimensional sphere of radius n converges in probability to the standard normal.

Results of a similar flavor can be found in convex geometry. It is well known that¹¹ as n grows large most of the volume of a ball is concentrated near any $n - 1$ dimensional hyperplane through the origin, so that the first coordinate of a uniform distribution over the ball is concentrated near the origin. While a uniform distribution over a ball is different than a uniform distribution over the sphere, in high dimensions the two are approximately the same.¹²

We note also that a rather different set of results apply if we replace the sphere with the surface of a hypercube. A hypercube in dimension n has $2n$ faces and projecting onto an axis collapses exactly two of them, so the projection of the uniform distribution with probability $1/(2n)$ is a point mass at each of the two endpoints and with probability $1 - 1/n$ is uniform over the line segment between those endpoints. Hence as $n \rightarrow \infty$ the probability mass does not collapse to the origin but approaches the uniform distribution. If the decision maker could choose directions randomly over the surface of a hypercube this would then imply that in high dimensions the optimal intervention would be approximately $\hat{\phi}^n = 1/2$. In order to properly orient the hypercube, however, the decision maker would have to know that the optimal solution lay along one of the coordinate axes. This shows that if the problem has additional structure that is known to the decision maker, the dimensionality n may not be a good measure of complexity.

More Directions

As we indicated earlier, the decision maker might have some ability to slightly improve the direction d . For example, the decision maker might be limited not

¹⁰See Diaconis and Friedman (1987) for a discussion of the history of the result and stronger versions.

¹¹See Ball (1997). We are grateful to an anonymous referee for this reference.

¹²To see this, observe that, since volume is of order R^n where R is the radius of the ball, the volume of a ball of radius $1 - \epsilon$ for large n is much smaller than the volume of a ball of radius 1. Hence most of the mass of a uniform over a ball in high dimensions is concentrated near the surface, which is to say, is approximately a uniform distribution over the boundary sphere.

just to the coordinate axis, but to directions that lie within an angle $\bar{\theta}$ of the coordinate axis, where $\bar{\theta}$ is relatively small. This corresponds to directions $\tilde{d}$ lying on the unit sphere with $|\tilde{d}'d| \geq \cos \bar{\theta}$. The most information that the decision-maker could have is to know the best direction within that cone. As long as $\bar{\theta} < \pi/2$ this does not effect the monotonicity result which did not depend on the range over which the integrals are taken, but rather on the fact that worse directions are more likely in higher dimensions.

Knowing the best direction in a cone does change the asymptotic results: with $\bar{\theta} = 0$ as in our base model, the optimal $\hat{\phi}(d)$ on the boundary where $d'\hat{x} = 0$ is $\hat{\phi}(d) = 0$ with corresponding loss $L(d) = 1$. With $\bar{\theta} > 0$ weight is still pushed towards the boundary as n increases, but now when d is orthogonal to $\hat{x}$ the direction of adjustment $\tilde{d}$ now has an angle $\pi/2 - \bar{\theta}/2$ lying closer to the optimum. This implies a boundary optimum of $\hat{\phi}(\tilde{d}) = \cos(\pi/2 - \bar{\theta}/2)$ with corresponding boundary loss of $1/2 - \cos(\pi/2 - \bar{\theta}/2)^2/2$. In practice, the decision maker is unlikely to know the best direction in the cone, and indeed our results argue it will be more difficult in higher dimensions. Hence these should be viewed as upper bounds on the size of the optimal intervention $\hat{\phi}$ and a lower bound on the loss L . The key point is that even for these bounds, if $\bar{\theta}$ is small the asymptotic intervention $\hat{\phi}(\tilde{d}) = \cos(\pi/2 - \bar{\theta}/2)$ is small, and the loss $1/2 - \cos(\pi/2 - \bar{\theta}/2)^2/2$ is close to $1/2$.

4. Quantitative Information and Gradient Descent

In the base model only qualitative information about the improvement is available: It is known only that in the direction d the loss is decreasing, but not how rapidly. We now consider the implications of quantitative information. Specifically, we ask what happens if the decision maker can move in a given direction d where the loss is decreasing, and in addition observe the derivative of the objective function in that direction. In this case the exact optimum in direction d can be obtained, though not necessarily the overall optimum because d will typically not be the best direction.¹³

For any given dimension n the quantitative information helps. However, as n grows, this advantage shrinks, and the loss converges to the same limit as in the base model without the quantitative information. Denote the loss in the direction d as $\lambda^n(d) = \min_{\phi} 1/2 - d_1\phi + (\phi)^2/2$, and let E^n denote the expectation operator with respect to the direction d , so that integrating over directions the expected loss is $E^n\lambda^n$.

Theorem 6. *The optimal solution $\hat{\phi}^n$ and corresponding loss λ^n satisfies $E^n\hat{\phi}^n = \hat{\phi}^n$, $L^{n+1} > E^{n+1}\lambda^{n+1} > E^n\lambda^n$ with $\lambda^1 = 0$ and $\lim_{n \rightarrow \infty} E^n\lambda^n = 1/2$.*

Proof. From the proof of Lemma 2 the objective function is $1/2 - d_1\phi + (\phi)^2/2$. Hence $\hat{\phi}^n(d) = d_1$. It follows that $E^n\hat{\phi}^n = \hat{\phi}^n = E^nd_1$, which is clear since the

¹³Note that if the decision maker could use knowledge of the derivative to choose a better direction this would lead to an even greater improvement, but we do not allow this.

best response is linear, so it does not matter if we first compute the optimum and then take the expectation as here, or first take the expectation and then compute the optimum as in the base model.

Substituting back into the objective function we have $\lambda^n(d) = 1/2 - d_1^2/2$ so that $E^n \lambda^n = 1/2 - E d_1^2/2$. Since $L^n = 1/2 - (E d_1)^2/2$ the result $L^{n+1} > E^{n+1} \lambda^{n+1}$ follows from Jensen's inequality. The remaining results follow from Lemma 5 which as observed in footnote 9 holds not only for computing $E^n d = E^n \cos \theta$ but also for computing $E^n d_1^2 = E^n (\cos \theta)^2$. $\square$

Gradient Descent

The availability of quantitative information about the direction d creates a link between our results and the method of gradient descent. We can index a problem of gradient descent by choosing a non-singular symmetric matrix A and a starting point y_0 with the objective function $(1/2)y' A' A y$. The first step of gradient descent computes the gradient $g = A y_0$ at y_0 , and then conducts a line search to find the optimum y_1 in along that line. The improvement made in this first step can be analyzed using our methods.

Recall that in our normalized coordinate system the origin has been moved to $x_0 = -(1, 0, 0, \dots, 0)^T$. Let B be a nonsingular linear transformation such that $B y_0 = x_0$ and $(A B^{-1})^T A B^{-1} = I$. In the coordinate system $x = B y$, the loss function is $(1/2)x^T x$, and the initial condition is x_0 . The direction of search is the transformed gradient $d = -B g$. Random choice of gradient descent problems, that is (A, y_0, B) then maps to random choices of d in our environment. This enables us to translate our results into conclusions about the first step of randomly chosen gradient descent problems. In particular, in the Appendix we show that if the probability distribution over gradient descent problems satisfies a symmetry and monotonicity condition then as the dimension of the problem increases the step size and gain from the first steps grow smaller and in the limit approach zero.

This result is about the first step of gradient descent algorithm, while the existing literature has focused on its rate of convergence. There is a standard convergence bound (see, for example, Meza (2010)) for the method. If the matrix A is randomly chosen by independently choosing the diagonal and upper triangle elements from a suitable distribution and filling in the rest by symmetry, a result on the distribution of eigenvalues by Wigner (1958) implies, as explained in the Appendix, that the convergence bound deteriorates as the dimension increases. Our result complements this by showing that as n increases not only does the algorithm perform more poorly in the limit of many iterations, but it performs more poorly in the first step as well.

Finally, we observe that when the decision maker observes the magnitude of the derivative, we can model decision makers who are less ignorant by supposing they can adjust the first m coordinates rather than just the first. We conjecture that similar results hold provided that m/n goes to 0.

5. Conclusion

We have defined and studied the “optimal intervention problem with limited information.” We modeled the complexity of a problem by its dimensionality, and found that as dimension increases it is optimal to make smaller interventions. Our proof follows the intuition that with more directions in which to move there is greater opportunity for mistake, and hence greater need for caution.

We consider two notions of limited information, both corresponding to what the decision maker knows about moving in any arbitrarily-chosen direction. In one formulation, only the sign of the derivative is known, and in the other its magnitude is known as well. In both cases we assume that the magnitude of the loss from non-intervention is known. In the “magnitude” case if line search is possible we show that this does not matter. More generally, while it may not be the case in practice that the magnitude of the loss is known, in more complex systems this information is likely to be more difficult to come by, which reinforces our results.

It is important to note that our result depends on the decision maker only observing the magnitude of the derivative in a fixed and small number of dimensions. In particular, if the decision maker observes the magnitude of the derivative in as few as n linearly independent directions then the full optimum is achieved.¹⁴

There is another point worth emphasizing. If $n > 1$ since $\hat{\phi}^n > 0$ it follows that while the expected loss is less than in the absence of intervention, there is still a positive probability that intervention will make the realized loss worse: Sometimes the intervention will make things worse rather than better.¹⁵

What do our theoretical results tell us? One case in point is that firms who do not respond to competition from complicated new products are sometimes criticized for doing nothing. However, if a firm lacks the technical expertise to locate the new “optimum,” and are limited to adjustments in a single dimension such as price, our results show that it may indeed be optimal to respond only slightly if at all. The responses of WordPerfect to the Windows 3.0 shock in 1990 and of Nokia and Blackberry to the iPhone shock in 2007 all fit into this category. In all three cases, firms did little to respond to the shock, and subsequent events showed that all indeed lacked the technical expertise to build competitive products.

One way to read our theory is that even in time of crisis large interventions are a bad idea. This is not the case. Instead, our model says that large ill-considered interventions are a bad idea. One example is that of the failure of NASDAQ computer systems during the Facebook IPO. Rather than study the

¹⁴This follows from the fact that the problem is quadratic, that the second derivatives are known, and that the distance to the optimum is known.

¹⁵This can be seen from the fact that if d is orthogonal to $\hat{x}$ the fact that $\hat{\phi}^n > 0$ implies a strictly greater loss than in the absence of intervention, hence the same must be true for d that are nearly orthogonal to $\hat{x}$.

system to understand the reason for the failure, programmers were instructed simply to make a large intervention in a single direction, namely to remove a validation check that had caused the system to shut down. The consequences were catastrophic: There was a cascading series of failures and “traders blamed NASDAQ for hundreds of millions of dollars of losses, and the mistake exposed the exchange to litigation, fines, and reputational costs.”¹⁶

References

- Ball, K. (1997). An elementary introduction to modern convex geometry. *Flavors of geometry*, 31, 1-58.
- Bellman, R. E., 1957. *Dynamic programming*. Princeton University Press, p.151.
- Clearfield, C. and Tilcsik A. 2018. How to prepare for a crisis you couldn't predict. *Harvard Business Review*, March 16.
- Cushman, F., Young, L. and Hauser, M., 2006. The role of conscious reasoning and intuition in moral judgment: Testing three principles of harm. *Psychological science*, 17(12), pp.1082-1089.
- Diaconis, P. and Friedman, D. (1987). A dozen de Finetti-style results in search of a theory. *Ann. Inst. Henri Poincaré: Probabilités et Statistiques*, Sup. au n° 2, 1987, Vol. 23, p. 397-423.
- Ely, J.C., 2011. Kludged. *American Economic Journal: Microeconomics*, 3(3), pp.210-31.
- Easterly, W., 2002. *The elusive quest for growth: economists' adventures and misadventures in the tropics*. MIT press.
- Meza, J.C., 2010. Steepest descent. *Wiley Interdisciplinary Reviews: Computational Statistics*, 2(6), pp.719-722.
- Samuelson, P.A., 1947. *Foundations of economic analysis*. Harvard University Press.
- Shulman, S., 2002. *Unlocking the sky*. Harper-Collins.
- Spranca, M., Minsk, E. and Baron, J., 1991. Omission and commission in judgment and choice. *Journal of experimental social psychology*, 27(1), pp.76-105.

¹⁶The story and the quote are from Clearfield and Tilcsik (2018), who also give examples where similar catastrophes were averted because complex malfunctioning systems were not restarted.

Stokey, N. L. 2009. *The Economics of Inaction: Stochastic Control models with fixed costs*. Princeton University Press, 2009.

Wigner, E., 1958. On the Distribution of the Roots of Certain Symmetric Matrices. *Annals of Mathematics* 67, 325-328.

Appendix: Gradient Descent

First Step of Gradient Descent

Recall that we have translated the coordinate system by subtracting the location of the optimum, so that the optimum $\hat{x}$ is at the origin and the status quo $x_0 = -(1, 0, 0, \dots, 0)^T$. We describe a gradient descent problem as a non-singular symmetric matrix A and a starting point y_0 with the objective function $(1/2)y^T A^T A y$. We map this into our setting with a nonsingular linear transformation B such that $B y_0 = x_0$ and $(AB^{-1})^T AB^{-1} = I$. The direction of gradient descent, that is the negative of the gradient $-g = A y_0$ maps to $d = B A y_0$.

Observe that

$$-x_0^T d = -x_0^T B A y_0 = -x_0^T B B' x_0 < 0$$

so that d is a direction of decrease for the objective function and lies on the half-sphere defined by x_0 and $\hat{x}$. If we find $\hat{\phi}^n(\mu)$ and $\lambda^n(\mu)$ in our problem then the optimum in the original problem is given by $\phi^n(\mu)g$ and the loss relative to $|y_0|^2$ is given by $\lambda^n(\mu)$.

Any distribution over (A, y_0, B) induces a distribution of d over our half-sphere. If the induced distribution over d is uniform describe the distribution over (A, y_0, B) as *pseudo-uniform*. Our main theorem then says that in this case as the dimension increases the expected size of the first step of gradient descent and expected gain from that step shrink to zero.

There are in general many pseudo-uniform choices of gradient descent problems. Here is one example: let e be uniform on our half-sphere and $A = \sqrt{2} \left((2/e_1) \left((e - (e_1 \hat{x})/2)(e - (e_1 \hat{x})/2)^T + (e_1^2/4)I \right) \right)^{1/2}$, $y_0 = A^{-1}x_0$, and $B = A$. Then as $\hat{x} = (1, 0, 0, \dots, 0)^T$

$$\begin{aligned} d &= 2A^{-1}(A^T)^{-1}\hat{x} \\ &= \left((2/e_1) \left((e - (e_1 \hat{x})/2)(e - (e_1 \hat{x})/2)^T + (e_1^2/4)I \right) \right) \hat{x} \\ &= (2/e_1) \left((e - (e_1 \hat{x})/2)(e - (e_1 \hat{x})/2)^T \hat{x} + (e_1^2/4)\hat{x} \right) \\ &= (2/e_1) \left((e - (e_1 \hat{x})/2)(e_1 - e_1/2) + (e_1^2/4)\hat{x} \right) \\ &= (2/e_1) \left((e_1/2)e - (e_1^2 \hat{x})/4 + (e_1^2/4)\hat{x} \right) = e \end{aligned}$$

is also uniform on our half-sphere so this distribution is indeed pseudo uniform. Note that since $B = A$ we can easily extend this to a measure with full support

over A . Take $y_0 = A^{-1}x_0$. The mapping $rd = BAy_0 = A^{-1}x_0$ induces equivalence classes of matrices A with two matrices being equivalent they give rise to the same value of r . Regardless of the distribution over matrices conditional on equivalent class, the example given provides a probability distribution over equivalence classes such that we get pseudo uniformity.

Note pseudo uniformity can be relaxed substantially. In proving Theorem 1 we conditioned on the angle θ . For given θ we used the fact that the distribution of directions d was symmetric. Let us say that a distribution over (A, y_0, B) is *symmetric* if this is the case. We view this as a fairly neutral assumption, that the gradient descent problem does not favor any particular direction. Second we used the fact that the distribution over angles $g^n(\theta)$ is uniform on $[0, \pi/2]$. However: the proof only requires that it be strictly positive and independent of n . Moreover, it is clear that if as we increase n the distribution for $n+1$ weakly stochastically dominates that for n , then this enhances the results given. If this is the case we say that the ensemble of distributions over (A, y_0, B) is *weakly monotone*. We conclude that symmetry plus weak monotonicity is sufficient for our results.

Wigner's Theorem and the Condition Number

A standard convergence bound (see, for example, Meza (2010)) shows that the rate of convergence declines as the ratio of the largest to smallest eigenvalue of $A'A$ grows. If the matrix A is randomly chosen by independently choosing the diagonal and upper triangle elements from a fixed and standardized distribution with moments of all orders, and filling in the rest by symmetry Wigner (1958)'s semi-circle law says that the fraction of normalized eigenvalues $\lambda_j/n^{1/2}$ of A that lie in an interval $[\lambda, \bar{\lambda}] \subseteq [-1, 1]$ converges in probability to

$$(2/\pi) \int_{\lambda}^{\bar{\lambda}} \sqrt{(1 - \lambda^2)} d\lambda.$$

Graphically on $[-1, 1]$ the function $\sqrt{(1 - \lambda^2)}$ is a semi-circle, hence the name. The implication for the ratio of the largest to smallest eigenvalue of AA' is clear: for any ϵ with high probability for large enough n there must be eigenvalues (and indeed quite a few of them) in $[0, 2\epsilon]$ and in $[1/2, 1/2 + 2\epsilon]$ so that the ratio is at least $1/\epsilon$. Hence, asymptotically, the convergence bound has no bite.