

1

Smoothed Analysis of Condition Numbers

Peter Bürgisser

*Institute of Mathematics**University of Paderborn**D-33098 Paderborn, Germany**email: pbuerg@upb.de***Abstract**

The running time of many iterative numerical algorithms is dominated by the condition number of the input, a quantity measuring the sensitivity of the solution with regard to small perturbations of the input. Examples are iterative methods of linear algebra, interior-point methods of linear and convex optimization, as well as homotopy methods for solving systems of polynomial equations. Thus a probabilistic analysis of these algorithms can be reduced to the analysis of the distribution of the condition number for a random input. This approach was elaborated upon for average-case complexity by many researchers.

The goal of this survey is to explain how average-case analysis can be naturally refined in the sense of smoothed analysis. The latter concept, introduced by Spielman and Teng in 2001, aims at showing that for all real inputs (even ill-posed ones), and all slight random perturbations of that input, it is unlikely that the running time will be large. A recent general result of Bürgisser, Cucker and Lotz (2008) gives smoothed analysis estimates for a variety of applications. Its proof boils down to local bounds on the volume of tubes around a real algebraic hypersurface in a sphere. This is achieved by bounding the integrals of absolute curvature of smooth hypersurfaces in terms of their degree via the principal kinematic formula of integral geometry and Bézout's theorem.

1.1 Introduction

In computer science, the most common theoretical approach to understanding the behaviour of algorithms is *worst-case analysis*. This means proving a bound on the worst possible performance an algorithm can

have. In many situations this gives satisfactory answers. However, there are cases of algorithms that perform exceedingly well in practice and still have a provably bad worst-case behaviour. A famous example is Dantzig's simplex algorithm. In an attempt to rectify this discrepancy, researchers have introduced the concept of *average-case analysis*, which means bounding the expected performance of an algorithm on random inputs. For the simplex algorithm, average-case analyses have been first given by Borgwardt (1982) and Smale (1983). However, while a proof of good average performance yields an indication of a good performance in practice, it can rarely explain it convincingly. The problem is that the results of an average-case analysis strongly depend on the distribution of the inputs, which is unknown, and usually assumed to be Gaussian for rendering the mathematical analysis feasible.

Spielman and Teng suggested in 2001 the concept of *smoothed analysis of algorithms*, which is a new form of analysis of algorithms that arguably blends the best of both worst-case and average-case. They used this new framework to give a more compelling explanation of the simplex method (for the shadow vertex pivot rule). For this work they were recently awarded the 2008 Gödel prize. See Spielman and Teng (2004) for the full paper.

The general idea of smoothed analysis is easy to explain. Let $T: \mathbb{R}^p \rightarrow \mathbb{R}_+ \cup \{\infty\}$ be any function (measuring running time etc). Instead of showing “it is unlikely that $T(a)$ will be large,” one shows that “for all $\bar{a}$ and all slight random perturbations $\bar{a} + \delta a$, it is unlikely that $T(\bar{a} + \delta a)$ will be large.” If we assume that the perturbation δa is centered (multivariate) standard normal with variance σ^2 , in short $\delta a \in N(0, \sigma^2)$, then the goal of a smoothed analysis of T is to give good estimates of

$$\sup_{\bar{a} \in \mathbb{R}^p} \text{Prob}_{\delta a \in N(0, \sigma^2)} \{T(\bar{a} + \delta a) \geq \epsilon^{-1}\}.$$

In a first approach, one may focus on expectations, that is on bounding

$$\sup_{\bar{a} \in \mathbb{R}^p} \mathbb{E}_{\delta a \in N(0, \sigma^2)} T(\bar{a} + \delta a).$$

Figure 1.1 succinctly summarizes the three types of analysis of algorithms.

Smoothed analysis is not only useful for analyzing the simplex algorithm, but can be applied to a wide variety of numerical algorithms. For doing so, understanding the concept of condition numbers is an important intermediate step.

A distinctive feature of the computations considered in numerical anal-

1. Smoothed Analysis of Condition Numbers 3

Worst-case analysis	Average-case analysis	Smoothed analysis
$\sup_{a \in \mathbb{R}^p} T(a)$	$\mathbb{E}_{a \in \mathcal{D}} T(a)$	$\sup_{\bar{a} \in \mathbb{R}^p} \mathbb{E}_{a \in N(\bar{a}, \sigma^2)} T(a)$

Fig. 1.1. Three types of analysis of algorithms. $\mathcal{D}$ denotes a probability distribution on $\mathbb{R}^p$.

ysis is that they are affected by errors. A main character in the understanding of the effects of these errors is the *condition number* of the input. This is a positive number which, roughly speaking, quantifies the errors when computations are performed with infinite precision but the input has been modified by a small perturbation. The condition number depends only on the data and the problem at hand. The best known condition number is that for matrix inversion and linear equation solving. For a square matrix A it takes the form $\kappa(A) = \|A\| \|A^{-1}\|$ and was independently introduced by Goldstine and von Neumann (1947) and Turing (1948).

Condition numbers are omnipresent in round-off analysis. They also appear as a parameter in complexity bounds for a variety of efficient iterative algorithms in linear algebra, linear and convex optimization, as well as homotopy methods for solving systems of polynomial equations. The running time $T(x, \epsilon)$ of these algorithms, measured as the number of arithmetic operations, can often be bounded in the form

$$T(x, \epsilon) \leq (\text{size}(x) + \mu(x) + \log \epsilon^{-1})^c, \tag{1.1}$$

with some universal constant $c > 0$. Here the input is a vector $x \in \mathbb{R}^n$ of real numbers, $\text{size}(x) = n$ is the dimension of x , the positive parameter ϵ measures the required accuracy, and $\mu(x)$ is some measure of conditioning of x . (Depending on the situation, $\mu(x)$ may be either a condition number or its logarithm. Moreover, $\log \epsilon^{-1}$ might be replaced by $\log \log \epsilon^{-1}$.)

We discuss the issue of condition-based analysis of algorithms in Sections 1.2–1.4, by elaborating a bit on the case of convex optimization and putting special focus on generalizations of Renegar’s (1995a, 1995b) condition number for linear programming. We also discuss Shub and Smale’s (1993a) condition number for polynomial equation solving.

Let us mention that L. Blum (1990) suggested to extend the complexity theory of real computation due to Blum, Shub, Smale (1989) by

measuring the performance of algorithms in terms of the size and the condition of inputs. However, up to now, no complexity theory over the reals has been developed that incorporates the concepts of approximation and conditioning and allows to speak about lower bounds or completeness results in that context.

Smale (1997) proposed a two-part scheme for dealing with complexity *upper bounds* in numerical analysis. The first part consists of establishing bounds of the form (1.1). The second part of the scheme is to analyze the distribution of $\mu(x)$ under the assumption that the inputs x are random with respect to some probability distribution. More specifically, we aim at tail estimates of the form

$$\text{Prob} \{ \mu(x) \geq \epsilon^{-1} \} \leq \text{size}(x)^c \epsilon^\alpha \quad (\epsilon > 0)$$

with universal constants $c, \alpha > 0$. In a first attempt, one may try to show upper bounds on the expectation of $\mu(x)$ (or $\log \mu(x)$, depending on the situation). Combining the two parts of the scheme, we arrive at upper bounds for the average running time of our specific numerical algorithms considered. So if we content ourselves with statements about the probabilistic average-case, we can eliminate the dependence on $\mu(x)$ in (1.1). This approach was elaborated upon for average-case complexity by Blum and Shub (1986), Renegar (1987), Demmel (1988), Kostlan (1988), Edelman (1988, 1992), Shub and Smale (1993b, 1994, 1996), Cheung and Cucker (2002), Cucker and Wschebor (2003), Cheung et al. (2005), Beltrán and Pardo (2007), and others. We only briefly discuss a few of these results in Section 1.5. Instead, we put emphasis on the analysis of the GCC-condition number $\mathcal{C}(A)$ of linear programming introduced by Goffin (1980) and Cheung and Cucker (2001), see (1.11). This is a variation of the condition number introduced by Renegar (1995a, 1995b). We discuss a recently found connection between the average-case analysis of the GCC-condition number and covering processes on spheres, and we present a sharp result on the probability distribution of $\mathcal{C}(A)$ for feasible inputs due to Bürgisser et al. (2007).

The main goal of this survey is to show that part two of Smale's scheme can be naturally refined by performing a smoothed analysis of the condition number $\mu(x)$ involved. This was already suggested by Spielman and Teng in their ICM 2002 paper. For the matrix condition number, results in this direction were obtained by Wschebor (2004) and Sankar et al. (2006). A recent paper by Tao and Vu (2007) deals with the matrix condition number under random discrete perturbations. Dunagan et al. (2003) gave a smoothed analysis of Renegar's condition

1. Smoothed Analysis of Condition Numbers

5

number of linear programming, thereby obtaining for the first time a smoothed analysis for the running time of interior-point methods, see also Spielman and Teng (2003).

A paper by Demmel (1988) has the remarkable feature that the probabilistic average-case analysis performed there for a variety of problems is not done with ad-hoc arguments adapted to the problem considered. Instead, these applications are all derived from a single result bounding the tail of the distribution of a conic condition number in terms of geometric invariants of the corresponding set of ill-posed inputs. Bürgisser et al. (2006, 2008) recently extended Demmel's result from average-case analysis to a natural geometric framework of smoothed analysis of conic condition numbers, called *uniform smoothed analysis*. This result will be presented in Section 1.6. The critical parameter entering these estimates turned out to be the degree of the defining equations of the set of ill-posed inputs. This result has a wide range of applications to linear and polynomial equation solving, as explained in Section 1.6.1. In particular, it easily gives a smoothed analysis of the condition number of a matrix. Moreover, Amelunxen and Bürgisser (2008) showed that this result, after suitable modification to a spherical convex setting, also allows a smoothed analysis of the GCC-condition number of linear programming.

The mathematical setting of uniform smoothed analysis has a clean and simple description. The set of ill-posed inputs to a computational problem is modelled as a subset Σ_S of a sphere S^p , which is considered the data space. In most of our applications, Σ_S is an algebraic hypersurface, but for optimization problems Σ_S will be semialgebraic. The corresponding conic condition number $\mathcal{C}(a)$ of an input $a \in S^p$ is defined as

$$\mathcal{C}(a) = \frac{1}{\sin d_S(a, \Sigma_S)},$$

where d_S refers to the angular distance on S^p . For $0 \leq \sigma \leq 1$ let $B(\bar{a}, \sigma)$ denote the spherical cap in the sphere S^p centered at $\bar{a} \in S^p$ and having angular radius $\arcsin \sigma$. Moreover, we define for $0 < \epsilon \leq 1$ the ϵ -neighborhood of Σ_S as

$$T(\Sigma_S, \epsilon) := \{a \in S^p \mid d_S(a, \Sigma_S) < \arcsin \epsilon\}.$$

The task of a uniform smoothed analysis of $\mathcal{C}$ consists of providing good upper bounds on

$$\sup_{\bar{a} \in S^p} \text{Prob}_{a \in B(\bar{a}, \sigma)} \{\mathcal{C}(a) \geq \epsilon^{-1}\},$$

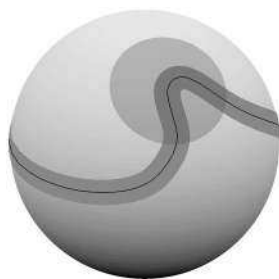


Fig. 1.2. Neighborhood of the curve Σ_S intersected with a spherical disk.

where a is assumed to be chosen uniformly at random in $B(\bar{a}, \sigma)$. The probability occurring here has an immediate geometric meaning:

$$\text{Prob}_{a \in B(\bar{a}, \sigma)} \{ \mathcal{C}(a) \geq \epsilon^{-1} \} = \frac{\text{vol}(T(\Sigma_S, \epsilon) \cap B(\bar{a}, \sigma))}{\text{vol}(B(\bar{a}, \sigma))}. \quad (1.2)$$

Thus uniform smoothed analysis means to provide bounds on the relative volume of the intersection of ϵ -neighborhoods of Σ_S with small spherical disks, see Figure 1.2. We note that uniform smoothed analysis interpolates transparently between worst-case and average-case analysis. Indeed, when $\sigma = 0$ we get worst-case analysis, while for $\sigma = 1$ we obtain average-case analysis. (Note that $S^p = B(\bar{a}, 1) \cup B(-\bar{a}, 1)$ for any $\bar{a}$.)

In Section 1.7 we explain the rich mathematical background behind our uniform smoothed analysis estimates. We first review classical results on the volume of tubes and then state the principal kinematic formula of integral geometry for spheres. Finally, in Section 1.7.3, we outline the proof of the main Theorem 1.2, which proceeds by estimating the integrals of absolute curvature arising in Weyl's tube formula (1939) with the help of Chern's (1966) principal kinematic formula and Bézout's theorem.

1.2 Condition Numbers for Linear Algebra

A numerical computation problem can often be formalized by a mapping $f: U \rightarrow Y$ between finite-dimensional real or complex vector spaces X and Y , where U is an open subset of X . The space X is interpreted as the set of inputs to the problem, Y is the set of solutions, and f is the solution map. Small perturbations δx of an input x result in a perturbation δy of the output $y = f(x)$. In order to quantify this effect

1. Smoothed Analysis of Condition Numbers

7

with regard to small relative errors, we choose norms on the spaces X and Y and define the *relative condition number* of f at x by

$$\kappa(f, x) := \lim_{\epsilon \rightarrow 0} \sup_{\|\delta x\| \leq \epsilon \|x\|} \frac{\|f(x + \delta x) - f(x)\| / \|f(x)\|}{\|\delta x\| / \|x\|}.$$

If f is differentiable at x , this can be expressed in terms of the operator norm of the Jacobian $Df(x)$ with respect to the chosen norms:

$$\kappa(f, x) = \|Df(x)\| \frac{\|x\|}{\|f(x)\|}.$$

In the case $X = Y = \mathbb{R}$, the logarithm of the condition number measures the *loss of precision* when evaluating f : if we know x up to ℓ decimal digits, then we know $f(x)$ roughly up to $\ell - \log_{10} \kappa(f, x)$ decimal digits.

Consider matrix inversion $f: \text{GL}(m, \mathbb{R}) \rightarrow \mathbb{R}^{m \times m}$, $A \mapsto A^{-1}$, measuring errors with respect to the L_2 -operator norm. A perturbation argument shows that the condition number of f at A equals the *classical matrix condition number*

$$\kappa(A) := \kappa(f, A) = \|A\| \|A^{-1}\|$$

of the matrix A . It is easy to see that $\kappa(A)$ also equals the condition number of the map $\text{GL}(m, \mathbb{R}) \rightarrow \mathbb{R}^m$, $A \mapsto A^{-1}b$ for fixed nonzero $b \in \mathbb{R}^m$. In fact, $\kappa(A)$ determines the condition number for solving a quadratic linear system of equations. It is also known that $\kappa(A)$ dominates the condition number of several other problems of numerical linear algebra, like the Cholesky and QR decomposition of matrices, see Amodei and Dedieu (2008). Moreover, the condition number $\kappa(A)$ appears in Wilkinson's round-off analysis of Gaussian elimination with partial pivoting (together with the so-called growth factor), see Wilkinson (1963) and Higham (1996).

Let us return to the problem of matrix inversion. We can interpret the set of singular matrices $\Sigma := \{A \in \mathbb{R}^{m \times m} \mid \det A = 0\}$ as its *set of ill-posed instances*. Let $\text{dist}(A, \Sigma)$ denote the distance of the matrix A to Σ , measured with respect to the L_2 -operator norm. The distance of A to Σ with respect to the Frobenius norm $\|A\|_F := (\sum_{i,j} a_{ij}^2)^{1/2}$ shall be denoted by $\text{dist}_F(A, \Sigma)$. The theorem of Eckart and Young (1936) states that $\text{dist}(A, \Sigma) = \text{dist}_F(A, \Sigma) = \|A^{-1}\|^{-1}$. As the right-hand side equals the smallest singular value of A , this is just a special case of the well-known fact that the k th largest singular value of A equals the distance of A to the set of matrices of rank less than k (with respect to the L_2 -operator norm). We rephrase Eckart and Young's result as the

following *condition number theorem*

$$\kappa(A) = \frac{\|A\|}{\text{dist}(A, \Sigma)} = \frac{\|A\|}{\text{dist}_F(A, \Sigma)}. \quad (1.3)$$

It is remarkable that the condition number $\kappa(A)$, which was defined using local properties, can be characterized in this global geometric way.

Demmel (1987) realized that this observation for the classical matrix condition number actually holds in much larger generality. For numerous computation problems, the condition number of an input x of norm one, say, can be bounded up to a constant factor by the inverse distance of x to a corresponding set of ill-posed inputs Σ . It is this key insight that allows to perform probabilistic analyses of condition numbers by geometric tools.

To further illustrate this connection, consider *eigenvalue computations*. Let $\lambda \in \mathbb{C}$ be a simple eigenvalue of $A \in \mathbb{C}^{m \times m}$. The sensitivity to compute λ from A is captured by a condition number $\kappa(A, \lambda)$, see (1.22). Wilkinson (1972) proved that

$$\kappa(A, \lambda) \leq \frac{\sqrt{2} \|A\|_F}{\text{dist}_F(A, \Sigma_{\text{eigen}})},$$

where Σ_{eigen} is the set of matrices in $\mathbb{C}^{m \times m}$ having a multiple eigenvalue.

Clearly, condition numbers are a crucial issue when dealing with finite precision computations and round-off errors. When considering iterative methods (instead of direct methods), it turns out that, even when assuming infinite precision arithmetic, the condition of an input often affects the number of iterations required to achieve a certain precision. A famous example for this phenomenon is the *conjugate gradient method* of Hestenes and Stiefel (1952). For a given linear system $Ax = b$, A a symmetric positive definite matrix, the conjugate gradient method starts with an initial value $x_0 \in \mathbb{R}^n$ and produces a sequence of iterates $x_0, x_1, \dots, x_n = x^*$ satisfying

$$\|x_k - x^*\|_A \leq 2 \left(\frac{\sqrt{\kappa(A)} - 1}{\sqrt{\kappa(A)} + 1} \right)^k \|x_0 - x^*\|_A,$$

where the A -norm of a vector v is defined as $\|v\|_A := (v^T A v)^{1/2}$. Therefore, roughly $\frac{1}{2} \sqrt{\kappa(A)} \ln \frac{1}{\epsilon}$ iterations are sufficient in order to achieve $\|x_k - x^*\|_A \leq \epsilon \|x_0 - x^*\|_A$.

1.3 Condition Numbers for Convex Optimization

We restrict our discussion to feasibility problems in convex conic form. Let X and Y be real finite-dimensional vector spaces endowed with norms. Further, let $K \subseteq X$ be a closed convex cone that is assumed to be regular, that is $K \cap (-K) = \{0\}$ and K has nonempty interior. We denote by $L(Y, X)$ the space of linear maps from Y to X endowed with the operator norm. Given $A \in L(Y, X)$, consider the feasibility problem of deciding

$$\exists y \in Y \setminus \{0\} \quad Ay \in K. \quad (1.4)$$

Two special cases of this general framework should be kept in mind. For $K = \mathbb{R}_+^n$, the nonnegative orthant in $\mathbb{R}^n$, one obtains the homogeneous *linear programming feasibility problem*. The feasibility version of homogeneous *semidefinite programming* corresponds to the cone $K = \mathcal{S}_+^n$ consisting of the positive semidefinite matrices in $\mathbb{R}^{n \times n}$.

The feasibility problem dual to (1.4) is

$$\exists x^* \in X^* \setminus \{0\} \quad A^*x^* = 0, \quad x^* \in K^*. \quad (1.5)$$

Here X^*, Y^* are the dual spaces of X, Y , respectively, $A^* \in L(X^*, Y^*)$ denotes the map adjoint to A , and $K^* := \{y^* \in Y^* \mid \forall x \in K \langle y^*, x \rangle \geq 0\}$ denotes the cone dual to K .

We denote by $\mathcal{D}$ the set of instances $A \in L(Y, X)$ for which the problem (1.4) is strictly feasible, i.e., there exists $y \in Y$ such that $Ay \in \text{int}(K)$. Likewise, we denote by $\mathcal{P}$ the set of $A \in L(Y, X)$ such that (1.5) is strictly feasible, i.e., there exists $x^* \in \text{int}(K^*)$ with $A^*x^* = 0$.

Both $\mathcal{D}$ and $\mathcal{P}$ are disjoint open subsets of $L(Y, X)$ and duality in convex optimization implies that $\mathcal{P}$ is the complement of the closure of $\overline{\mathcal{D}}$ in $L(Y, X)$, cf. Boyd and Vandenberghe (2004). The *conic feasibility problem* is to decide for given $A \in L(Y, X)$ whether (1.4) or (1.5) holds. The common boundary $\Sigma := \partial\mathcal{D} = \partial\mathcal{P}$ of the sets $\mathcal{D}$ and $\mathcal{P}$ can be considered as the set of ill-posed instances. Indeed, for given $A \in \Sigma$, arbitrarily small perturbations of A may yield instances in both $\mathcal{D}$ and $\mathcal{P}$.

Renegar (1995a) defined the *condition number* of the conic feasibility problem by

$$C(A) := \frac{\|A\|}{\text{dist}(A, \Sigma)}. \quad (1.6)$$

He observed that the number of steps of interior-point algorithms solving the conic feasibility problem can be effectively bounded in terms

of $C(A)$. Before elaborating on this important issue, let us characterize the condition number $C(A)$ in a different way. Suppose there exists $e \in \text{int}(K)$ such that the unit ball $B(e, 1)$ centered at e is contained in K . We define $\lambda_{\min}: X \rightarrow \mathbb{R}$ by $\lambda_{\min}(x) := \max\{t \in \mathbb{R} \mid x - te \in K\}$ and note that $x \in K \Leftrightarrow \lambda_{\min}(x) \geq 0$. For $K = \mathbb{R}_+^n$ and $e = (1, \dots, 1)$ we have $\lambda_{\min}(x) = \min_i x_i$, while in the case $K = \mathcal{S}_+^n$ and e being the unit matrix, $\lambda_{\min}(x)$ equals the minimum eigenvalue of x .

The problem (1.4) is feasible iff there exists $y \in Y$ of norm one such that $\lambda_{\min}(Ay) \geq 0$. A vector y maximizing $\lambda_{\min}(Ay)/\|y\|$ may be interpreted as a best-conditioned solution, due to the following max-min characterization in Cheung et al. (2008):

$$\text{dist}(A, \Sigma) = \left| \max_{\|y\|=1} \lambda_{\min}(Ay) \right|. \quad (1.7)$$

Actually, in that paper a more general result is shown. Suppose we have a multifold conic structure: $X = X_1 \times \dots \times X_r$, where X_i is a normed vector space, $K = K_1 \times \dots \times K_r$ with regular closed convex cones K_i in X_i , and $e_i \in \text{int}(K_i)$ such that the unit ball centered at e_i is contained in K_i . We have a corresponding function $\lambda_{\min}^i: X_i \rightarrow \mathbb{R}$. Then (1.4) can be written as

$$\exists y \in Y \setminus \{0\} \quad A_1 y_1 \in K_1, \dots, A_r y_r \in K_r,$$

where $A_i \in L(Y, X_i)$ is the composition of A with the projection onto X_i . Generalizing (1.6), we define the corresponding *multifold condition number* $\mathcal{C}(A)$ by

$$\mathcal{C}(A) := \left(\min_{B \in \Sigma} \max_i \frac{\|A_i - B_i\|}{\|A_i\|} \right)^{-1}.$$

It is easy to see that $\mathcal{C}(A) \leq C(A)$ when taking $\|A\| = \max_i \|A_i\|$. Note that in the case $r = 1$ of just one factor, we retrieve $C(A) = \mathcal{C}(A)$. The condition number $\mathcal{C}(A)$ seems a more natural measure of conditioning in the multifold setting, when allowing component normalization as preconditioning. Cheung et al. (2008) proved the following *condition number theorem*, extending (1.7),

$$\frac{1}{\mathcal{C}(A)} = \left| \max_{\|y\|=1} \min_i \frac{\lambda_{\min}^i(A_i y)}{\|A_i\|} \right|. \quad (1.8)$$

Let us now have a closer look at the important special case of $X_i = \mathbb{R}$, $K_i = \mathbb{R}_+$, $e_i = 1$ for $i = 1, \dots, n$. We endow $X = \mathbb{R}^n$ with the L_∞ -norm and $Y := \mathbb{R}^{m+1}$ with the L_2 -norm. The problem (1.4) now reads as the