

Lecture #11: Discrete Probability for Computer Science II

Key Concepts

Tail Bounds

Basic Bounds

Theorem (Basic Inequality). Let Ω be a **finite** sample space with probability distribution $P : \Omega \rightarrow \mathbb{R}$, let $X : \Omega \rightarrow \mathbb{R}$ be a random variable, and let $h : \mathbb{R} \rightarrow \mathbb{R}$ be a total function such that

$$h(x) \geq 0 \quad \text{for all } x \in \mathbb{R}.$$

Then, for every real number a such that $a > 0$,

$$P(h(X) \geq a) \leq \frac{E[h(X)]}{a}.$$

Corollary (Markov's Inequality). Let Ω be a **finite** sample space with probability distribution $P : \Omega \rightarrow \mathbb{R}$, and let $X : \Omega \rightarrow \mathbb{R}$ be a random variable. Then, for every **positive** real number a ,

$$P(|X| \geq a) \leq \frac{E[|X|]}{a}.$$

Variance and Standard Deviation

Definition. Let Ω be a sample space with probability distribution $P : \Omega \rightarrow \mathbb{R}$, and let $X : \Omega \rightarrow \mathbb{R}$. Then the **variance** of X , with respect to P , is

$$\text{var}(X) = \sum_{\mu \in \Omega} (X(\mu) - E[X])^2 \times P(\mu)$$

and the **standard deviation** of X , denoted $\sigma(X)$, is $\sqrt{\text{var}(X)}$.

Theorem. Let Ω be a **finite** sample space, let $P : \Omega \rightarrow \mathbb{R}$ be a probability distribution for Ω , and let X be a random variable. Then X^2 is also a random variable, and

$$\text{var}(X) = E[X^2] - E[X]^2.$$

It is *not* generally true that $\text{var}(X + Y) = \text{var}(X) + \text{var}(Y)$ for a pair of random variables X and Y — so the following result is useful (and not trivial):

Theorem. Let Ω be a **finite** sample space with probability distribution $P : \Omega \rightarrow \mathbb{R}$ and let $X_1, X_2, \dots, X_n : \Omega \rightarrow \mathbb{R}$ be random variables (for some positive integer n). If $X_1, X_2, \dots, X_n$ are **pairwise independent** then

$$\text{var}(X_1 + X_2 + \dots + X_n) = \text{var}(X_1) + \text{var}(X_2) + \dots + \text{var}(X_n).$$

More Bounds

Theorem (Chebyshev's Inequality). Let Ω be a **finite** sample space with probability distribution $P : \Omega \rightarrow \mathbb{R}$, let X be a random variable, and let $a \in \mathbb{R}$ such that $a > 0$. Then

$$P(|X| \geq a) \leq \frac{E[X^2]}{a^2}.$$

and

$$P(|X - E[X]| \geq a) \leq \frac{\text{var}(X)}{a^2}.$$

Theorem (Cantelli's Inequality). Let Σ be a **finite** sample space with probability distribution $P : \Omega \rightarrow \mathbb{R}$, let $X : \Omega \rightarrow \mathbb{R}$ be a random variable, and let $a \in \mathbb{R}$ such that $a > 0$. Then

$$P(X - E[X] \geq a) \leq \frac{\text{var}(X)}{a^2 + \text{var}(X)}.$$

Cantelli's Inequality is sometimes called the “One-Sided Chebyshev's Inequality”.

Theorem (Chernoff Bound). Let Ω be a finite sample space with probability distribution $P : \Omega \rightarrow \mathbb{R}$. Suppose that $X_1, X_2, \dots, X_n$ are mutually independent random variables such that $X_i : \Omega \rightarrow \{0, 1\}$ for $1 \leq i \leq n$, and suppose that $P(X_i = 1) = p$ for every integer i such that $1 \leq i \leq n$, for a real number p such that $0 \leq p \leq 1$. Let $X = X_1 + X_2 + \dots + X_n$. Then, for every real number θ such that $0 \leq \theta \leq 1$,

$$P(X \geq (1 + \theta)pn) \leq e^{-\frac{\theta^2}{3}pn}.$$

Applications to Computer Science: Design and Analysis of Algorithms

Average-Case Analysis of Deterministic Algorithms

Consider a deterministic algorithm that solves a given computational problem. Let n be a non-negative integer.

- A **sample space** consists of instances of this problem of size n — or of *sets* of these inputs (such that the inputs in any one of these sets would cause the algorithm to behave in the same way).
- A **probability distribution** is formed, to model assumptions about the *likelihood* of the (types of) inputs included in the sample space.
- The number of steps used by the algorithm, on each (type of) input in the sample space, is used to define a **random variable**.
- The **expected running time** of the algorithm, for input size n , is the **expected value** of this random variable with respect to this probability distribution.

Examples from CPSC 331: Analysis of algorithms for operations in hash tables; analysis of the deterministic Quicksort algorithm.

Randomized Algorithms

Randomized algorithms use random number generators during the execution — so that neither the output generated, nor the number of steps used to generate it, are necessarily *fixed*, even when the algorithm's input is.

Analysis for a Fixed Input

Consider the execution(s) of a randomized algorithm, for a given computational problem, on a given instance of this problem.

- A **sample space** consists of sequences values that are produced during the algorithm's execution, using the random number generator.

- A **probability distribution** is formed, to model assumption about the random number generator that is being used. It is usually assumed that this is a “perfect” source of random values: When a value from a finite, nonempty set is requested then each element of the set is produced with the same probability, and this is independent of all choices that have been made before this (or anything else).
- The number steps that are used by the algorithm, for this input and for this sequence of random variables, is used to define a **random variable**.
- The **expected running time**, for this input, is the expected value of this random variable, with respect to this probability distribution.

Measuring Running Time as a Function of Input Size

- Suppose that the *expected running time* of a randomized algorithm, for a given instance of the computational problem, is as defined above. Then the **worst-case running time** of this algorithm (with respect to this probability distribution) is a function $T : \mathbb{N} \rightarrow \mathbb{R}$.

If there are finitely many instances of the problems with size n — or only finitely many values that that the “expected running time” can assume, for instances of the problem with size n , then the worst-case expected running time, $T(n)$, is the *maximum* of the expected running times of the algorithm for all inputs of size n .¹

Example from CPSC 331: Analysis of a randomized Quicksort algorithm.

Randomized Algorithms That Can Fail

A **decision problem** is a computational problem that answers a “Yes-or-No” question, so that its output is always either `true` (corresponding to the answer “Yes”) or `false` (corresponding to the answer “No”).

- A **Las Vegas** algorithm for a decision problem is a randomized algorithm that can never give a wrong answer. Las Vegas algorithms are of interest when their worst-case expected running times are at most polynomial in the size of the algorithm’s input.
- A **Monte Carlo** algorithm for a decision problem is a randomized algorithm that can, sometimes, fail — but that does so with small probability:

¹This to be defined to be the *supremum* of this set of values, when infinitely many inputs must be considered.

- The algorithm only returns true, when executed on a given input, if this is the correct answer for that input. That is, the algorithm always returns the corrected answer when it is executed on an input for which correct answer is “No”.
- If the algorithm is executed on an input for which the correct answer is “Yes” (so that the algorithm *should* return true) then the probability that the algorithm *does not* return true is at most $\frac{1}{2}$.
- Other kinds of algorithms that can fail are studied in the literature too. This will be considered, further, in the tutorial exercise for this topic.

What Really Happens

Randomized algorithms are almost always analyzed using the assumption that “perfect” sources of random values are available — that is, when a value is to be chosen from a finite non-empty set then every value is returned with the same probability, and this is independent of all previous uses of the random number generator.

- Current programming environments currently provide **pseudorandom number generators** instead. These are *deterministic* processes that generate sequences of values, starting from an initial value (generally called the “seed”).
- Thus the *assumption*, used to define probability distributions in order to analyze the performance of randomized algorithms, is **false**.
- Nevertheless, experiments (involving monitoring the performance of randomized algorithms, as they are used in progress) generally given results that are consistent with the results that these flawed analyses indicate.
- A huge amount of additional information about this is available — and this is beyond the scope of this course.