

St. Francis Xavier University
Department of Computer Science
CSCI 541: Theory of Computing
Lecture 6: Probabilistic Computation
Winter 2026

When we defined and compared deterministic and nondeterministic Turing machines, we noted that using nondeterminism in a computation could provide us with desirable behaviour; instead of deterministically following the same steps and ending up with the same outcomes, we could “make guesses” at certain branching points in the computation and see what outcomes arise as a result. By running all of these computation branches in parallel, we can evaluate every possibility at once to see if any branch produces an accepting computation.

However, despite how appealing this behaviour sounds, it doesn’t actually give us any additional power to compute more things. Any nondeterministic computation can be simulated by a deterministic Turing machine with some penalty incurred in the amount of resources used. Indeed, since real-world computers behave in an entirely deterministic manner, we can’t implement *true* nondeterminism in any algorithm we write—it is, in effect, an academic notion that we use to reason about aspects of computability and complexity.

Since nondeterministic behaviour is desirable yet we can’t implement it directly, we must settle for some approximation of nondeterminism, and this is where *probabilistic computation* enters. Probabilistic machines are effectively an intermediary: they are deterministic machines that simulate nondeterministic choices by way of consulting random values. Although a probabilistic machine can’t evaluate every computation branch at once in the way that a nondeterministic machine can, it can still “make guesses” by following one of potentially many computation branches with some prespecified odds.

Of course, this presents a downside: if the “guesses” made by a probabilistic machine are incorrect, this may lead to the machine producing an incorrect output, say by mistakenly rejecting a word that does in fact belong to the language of the machine. This naturally isn’t an issue with nondeterministic machines, since all computation branches are evaluated at once and we only require the existence of one accepting computation branch to produce a positive outcome. But despite the risk of incorrect outputs, probabilistic machines do possess remarkable utility: by allowing the machine to produce incorrect outputs with a small probability, we are often able to solve problems in less time or space than would be used by a deterministic machine.

1 Probabilistic Turing Machines

In order for us to perform probabilistic computations, we must define a model of computation that is capable of using randomness. This model of computation, a *probabilistic Turing machine*, is fundamentally the same as any other Turing machine we’ve seen thus far, but its transition “function” is split into two distinct functions δ_1 and δ_2 that the machine applies with equal probability. That is, the machine applies δ_1 with probability $1/2$ or otherwise it applies δ_2 , akin to flipping a fair coin.

Definition 1 (Probabilistic Turing machine). A probabilistic Turing machine is a tuple $(Q, \Sigma, \Gamma, \delta_1, \delta_2, q_0, q_{\text{accept}}, q_{\text{reject}})$, where

- Q is a finite set of *states*;
- Σ is the *input alphabet* (where $\sqcup \notin \Sigma$);
- Γ is the *tape alphabet* (where $\sqcup \in \Gamma$ and $\Sigma \subseteq \Gamma$);
- $\delta_1 : (Q \setminus \{q_{\text{accept}}, q_{\text{reject}}\}) \times \Gamma \rightarrow Q \times \Gamma \times \{L, R\}$ is the *first transition function*;

- $\delta_2 : (Q \setminus \{q_{\text{accept}}, q_{\text{reject}}\}) \times \Gamma \rightarrow Q \times \Gamma \times \{L, R\}$ is the *second transition function*;
- $q_0 \in Q$ is the *initial* or *start state*;
- $q_{\text{accept}} \in Q$ is the *final* or *accepting state*; and
- $q_{\text{reject}} \in Q$ is the *rejecting state*.

Note that δ_1 and δ_2 are fundamentally the same function in that they each map a tuple of state and tape symbol to a tuple of state, tape symbol, and input head movement, but they do not necessarily have the same behaviour. As an illustrative example, supposing we are in some state q_i and reading some symbol a , δ_1 may produce the tuple (q_j, b, R) while δ_2 may produce the tuple (q_k, b, L) . In essence, we have two choices for each transition, and each choice occurs with probability $1/2$.

Remark. We may alternatively define a probabilistic Turing machine to be a deterministic Turing machine that has access to a special read-only *random tape* filled with random bits. Under this definition, the machine can choose which of its transition functions to apply by consulting the bits on the random tape.

Computations and Accepting Computations. While a deterministic computation either always accepts or always rejects its input word, and a nondeterministic computation accepts if there exists at least one accepting computation branch, probabilistic computations behave in a somewhat different manner. To each input word, a probabilistic Turing machine associates a value between 0 and 1 that corresponds to the probability that a randomly selected computation branch accepts that word. Thus, some percentage of computation branches will accept the word, while the other branches will reject it.

Since our definition specifies that a probabilistic Turing machine applies either transition function with probability $1/2$, every step of the machine's computation can proceed in one of two ways, and so the computation tree resembles a binary tree. For deterministic portions of the computation, both branches will lead to the same configuration, while for nondeterministic portions of the computation having more than two possible outcomes, the structure of the computation can be reorganized to produce exactly two branches at each step.

In this way, for any probabilistic Turing machine $\mathcal{M}$, we can assign to each branch b of its computation tree the probability

$$\mathbb{P}[b] = \frac{1}{2^k},$$

where k is the number of computation steps that occurred along the branch b . As an immediate consequence, the probability that $\mathcal{M}$ accepts an input word w is

$$\mathbb{P}[\mathcal{M} \text{ accepts } w] = \sum_{b \in \text{ACC}} \mathbb{P}[b],$$

where ACC is the set of all halting computations on w that lead to an accepting state of $\mathcal{M}$. At the same time, the probability that $\mathcal{M}$ rejects w is simply

$$\mathbb{P}[\mathcal{M} \text{ rejects } w] = 1 - \mathbb{P}[\mathcal{M} \text{ accepts } w].$$

Lastly, the language of a probabilistic Turing machine is the set of all words whose probability of acceptance is above some predefined threshold. There are many ways we can define such a threshold, and these myriad ways form the cornerstone of the study of randomized complexity theory.

2 Probabilistic Resource-Bounded Computations

When it comes to measuring the resource usage of a probabilistic Turing machine, things aren't quite as straightforward as when we were dealing with deterministic or nondeterministic Turing machines. Since a probabilistic Turing machine chooses computation branches at random, we have no way of knowing in advance whether we might follow a branch that halts quickly or a branch that takes ages to complete. Thus,

before we can define time and space bounds for probabilistic Turing machines, we must come up with some way to accurately measure the resources used by an individual computation branch that is chosen at random!

To make our job easier, we will introduce the following assumption on the structure of the computation tree of any probabilistic Turing machine computing a length- n input word in time at most $t(n)$.

ABEL property. Let $\mathcal{M}$ be a $t(n)$ -time probabilistic Turing machine. The computation of $\mathcal{M}$ on a length- n input word and with any sequence of random bits (or “coin tosses”) halts in exactly $t(n)$ steps.

The acronym ABEL stands for “all branches, equal length”; under this assumption, every branch of the computation tree has the same length, which makes our job of analyzing the computation much easier.

If we suppose the computation of some probabilistic Turing machine $\mathcal{M}$ satisfies the ABEL property, then since every computation step of $\mathcal{M}$ can lead to one of two choices and since every computation of $\mathcal{M}$ takes time $t(|w|)$ by the ABEL property, the computation tree of $\mathcal{M}$ consists of exactly $2^{t(|w|)}$ branches. For any input word w , the probability that $\mathcal{M}$ accepts w is given by the expression

$$\mathbb{P}[\mathcal{M} \text{ accepts } w] = \frac{|\text{ACC}|}{2^{t(|w|)}},$$

where ACC again denotes the set of all halting computations on w that lead to an accepting state of $\mathcal{M}$.

Thankfully, the ABEL property is not so strong that we lose generality by taking it as an assumption. As long as $t(n)$ is a time-constructible function, we can take any arbitrary $t(n)$ -time probabilistic Turing machine and construct an equivalent probabilistic Turing machine that satisfies the ABEL property.

3 Las Vegas and Monte Carlo Algorithms

Due to their inherent computational behaviour, probabilistic Turing machines run *randomized algorithms*. However, when it comes to an algorithm, not all randomized behaviour is the same: some algorithms differ in how many resources are used over the course of the computation, while others differ in the probability that the algorithm outputs a correct answer.

We can classify all randomized algorithms as one of two types (well, strictly speaking, three types), depending on whether the randomization aspect applies to the runtime of the algorithm or to the output produced by the algorithm.

- A *Las Vegas algorithm* always gives us a correct output, but the runtime of the algorithm may vary depending on the random choices made by the algorithm.
- A *Monte Carlo algorithm* is guaranteed to run in a certain amount of time, but the output produced may be subject to error depending on the random choices made by the algorithm. Specifically:
 - A Monte Carlo algorithm *with one-sided error* is always correct when it returns one answer (“yes”) and is incorrect with some bounded probability when it returns the other answer (“no”), or vice versa.
 - A Monte Carlo algorithm *with two-sided error* is incorrect with some bounded probability regardless of the answer it returns (“yes” or “no”).

As a small example, suppose we have an array A containing n elements: $(n - 1)$ zeroes and 1 one. We can design two randomized algorithms to find the index value i such that $A[i] = 1$:

Algorithm 1: Array search—Las Vegas

```

while true do
     $i \leftarrow$  random integer
    if  $A[i] = 1$  then
        return  $i$ 
    
```

Algorithm 2: Array search—Monte Carlo

```

for  $0 \leq c \leq 10$  do
     $i \leftarrow$  random integer
    if  $A[i] = 1$  then
        return  $i$ 
    return failure
    
```

With the Las Vegas algorithm, we're guaranteed to find the index value, but we could potentially loop for a very long time if we keep making unlucky random choices. With the Monte Carlo algorithm, we only loop 10 times; we will always return the index value if it is found, but we may instead return "failure" if we don't find the index value (despite the fact that it exists). Therefore, this particular Monte Carlo algorithm has one-sided error.

The relationship between Las Vegas and Monte Carlo algorithms can be summarized in the following table:

	Correctness	Runtime
Las Vegas	Certain	Uncertain
Monte Carlo	Uncertain	Certain

Alternatively, if you prefer a mnemonic to distinguish between the two, remember that Las Vegas algorithms have a Variable runtime, while Monte Carlo algorithms sometimes make Mistakes.

Having distinguished between the two types of randomized algorithms, a natural question to ask might be "can we convert between the two?" That is, if we have a randomized algorithm of one type, can we adapt the way it applies randomization to be of the other type? The answer is... sometimes.

If we start with a Las Vegas algorithm, it's rather straightforward for us to adapt the algorithm to run in a fixed amount of time at the expense of introducing error into the output we get.

Theorem 2. *Every Las Vegas algorithm can be converted to a Monte Carlo algorithm.*

Proof. Omitted. □

On the other hand, we cannot in general convert a Monte Carlo algorithm to a Las Vegas algorithm unless we have some method of testing the correctness of the output produced by the algorithm. If we have such a testing method, then we can simply repeat the execution of the Monte Carlo algorithm until we get a correct output. Otherwise, we must rely on whatever knowledge we have of the distribution of outputs in order to determine our level of confidence in the output given by a Monte Carlo algorithm.

4 Randomized Time Complexity Classes

As with our deterministic and nondeterministic algorithms, there is a rich tapestry of complexity classes that characterize the behaviour of randomized algorithms under different conditions. Depending on the type of randomized algorithm we're taking under consideration, and on the particular threshold we fix on the probability of acceptance, we obtain one of many randomized complexity classes. Each of these complexity classes also interacts in interesting ways with our fundamental classes like P, NP, and the rest. In this section, we consider the major randomized complexity classes in turn.

4.1 Two-Sided Error: PP

We will begin our study of randomized complexity theory by devising a probabilistic analogue of the class P; that is, a class comprised of decision problems that can be solved by a probabilistic Turing machine in some polynomial amount of time $p(n)$, independent of any random choices made over the course of the computation. Supposing that $p(n)$ is time constructible, it's easy for the probabilistic Turing machine to abide by this time bound: it simply counts the number of computation steps it performs and halts once the limit has been reached.

The simplest way we can define a randomized complexity class is just to fix some threshold value and state that an input word will be accepted by the probabilistic Turing machine if its probability of acceptance is greater than this threshold.

Since the problems we are currently considering can be solved by a polynomial-time probabilistic Turing machine, the name of our first complexity class will be PP, representing all *probabilistic polynomial time* decision problems.

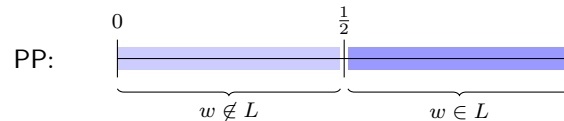
Definition 3 (The class PP). A decision problem L belongs to the complexity class PP if there exists a probabilistic Turing machine $\mathcal{M}$ such that, given an input word w ,

- if $w \in L$, then $\mathbb{P}[\mathcal{M} \text{ accepts } w] > 1/2$; and
- if $w \notin L$, then $\mathbb{P}[\mathcal{M} \text{ accepts } w] < 1/2$.

Note that the class PP corresponds to the class of decision problems that can be solved by a Monte Carlo algorithm with two-sided error. If $w \in L$, then a PP-machine will reject w with probability less than $1/2$, while if $w \notin L$, then the machine will accept w with probability less than $1/2$.

For those who are curious why we selected $1/2$ as our threshold, there's no particular reason why we selected that value specifically. Indeed, we could choose any value $0 < \alpha < 1$ as our threshold, and the definition of PP would follow in exactly the same way for α as it did for $1/2$. We simply went with $1/2$ as it's a "nice" value to work with.

Now, with randomized complexity classes, it's often helpful to illustrate on a probability line where exactly the accepting and rejecting probabilities may fall. Probability lines are a great tool to compare and contrast randomized complexity classes, since they allow us to differentiate classes at a glance. If we were to draw a "picture" of the class PP on a probability line and highlight the regions of this line where the probability of acceptance would appear given either $w \in L$ or $w \notin L$, we would have the following:



Unfortunately, as you may have noticed in our illustration, one consequence of our definition of the class PP is that there exists a tricky *probability gap* between the probability of a PP-machine accepting inputs $w \in L$ and the probability of the same machine accepting inputs $w \notin L$. Since the only condition is that we accept words $w \in L$ with probability strictly greater than $1/2$, it is possible to squeeze the acceptance probability to be arbitrarily close to $1/2$ and make it increasingly difficult for us to distinguish between correct and incorrect outputs.

For example, it is possible for us to design a PP-machine that accepts inputs $w \in L$ with probability $1/2 + 1/2^n$ and accepts inputs $w \notin L$ with probability $1/2 - 1/2^n$, where $n = |w|$. As n grows, the gap between correctly accepting a word in L and incorrectly accepting a word not in L becomes arbitrarily small, and we must devote increasing attention to making sure our outputs are correct.

While the existence of an arbitrarily small probability gap alone suggests that working with the class PP can be difficult, we can go one step further by showing that PP is such a powerful class that it contains problems that are not known to be efficiently computable!

Recalling the class NP, we can define this class probabilistically as the class of all decision problems L for which there exists a probabilistic Turing machine $\mathcal{M}$ where, for all $w \in L$, $\mathbb{P}[\mathcal{M} \text{ accepts } w] > 0$. This aligns with the definition of a nondeterministic computation, where we accept an input word if there exists at least one accepting path in the computation tree.

Using our probabilistic definition of the class NP, we can prove the following result.

Theorem 4. $\text{NP} \subseteq \text{PP} \subseteq \text{PSPACE}$.

Proof. First, we show that $\text{NP} \subseteq \text{PP}$. Consider the computation tree of any nondeterministic Turing machine $\mathcal{M}$ recognizing a decision problem L . We can construct a PP-machine $\mathcal{M}'$ recognizing the same decision problem L by taking $\mathcal{M}$, adding a new initial state, and adding a nondeterministic transition out of the new initial state to one of two subtrees: the first subtree is rooted at the original computation tree of $\mathcal{M}$, while the second subtree is simply a leaf corresponding to an accepting computation. The computation tree of $\mathcal{M}'$ is depicted in Figure 1.

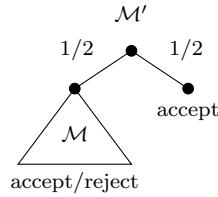


Figure 1: The computation tree of the PP-machine $\mathcal{M}'$ constructed in the proof of Theorem 4.

Since $L \in \text{NP}$, if $w \in L$, then there exists at least one accepting computation in the computation tree of $\mathcal{M}$, so $\mathbb{P}[\mathcal{M} \text{ accepts } w] > 0$ following the probabilistic definition of NP. As a consequence, in the computation tree of $\mathcal{M}'$, we accept with probability $1/2$ if we enter the “accepting computation” subtree from the initial configuration, and we accept with probability greater than 0 if we enter the “computation tree of $\mathcal{M}$ ” subtree from the initial configuration. Therefore, if $w \in L$, then $\mathbb{P}[\mathcal{M}' \text{ accepts } w] > 1/2$. We can make a similar argument when $w \notin L$. Thus, $\mathcal{M}'$ is a PP-machine.

To see that $\text{PP} \subseteq \text{PSPACE}$, observe that the height of the computation tree of a PP-machine is bounded by some polynomial, since every computation branch runs in polynomial time. Thus, we can check every possible computation branch using a polynomial amount of space and count the number of accepting branches using a logarithmic-space binary counter. Overall, this process requires a polynomial amount of space. $\square$

Remark. You may have noticed in our proof showing $\text{NP} \subseteq \text{PP}$ that, if $w \notin L$, then our computation tree for $\mathcal{M}'$ accepts w with probability exactly $1/2$, while our definition of PP states that we must accept w with probability strictly less than $1/2$! Fortunately, this doesn’t present a problem due to our earlier observation that we can select any value $0 < \alpha < 1$ as our threshold and obtain a reasonable definition of the class PP.

Lastly, it is known that the class PP is closed under complement. Thus, by adapting our previous proof, we can similarly show with not too much additional work that $\text{coNP} \subseteq \text{PP}$.

4.2 Two-Sided Bounded Error: BPP

In light of the fact that PP contains NP, it is clear that working with PP directly may be more trouble than it’s worth, given that its definition is so general. What would be nice, then, is to come up with a more refined complexity class that “acts like” PP, but avoids the pitfall of arbitrarily small probability gaps.

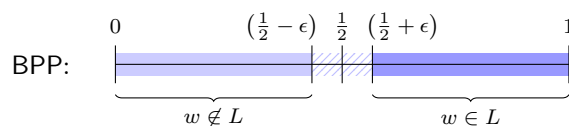
We can achieve this by enforcing the property that the probability gap maintains a minimum size; in other words, that our probabilistic Turing machine accepts inputs $w \in L$ with some probability $1/2$ plus a constant. This means that the probability gap cannot become arbitrarily small by, say, relying on the size of the input. By bounding the minimum size of the probability gap, we obtain the class BPP of *bounded probabilistic polynomial time* decision problems.

Definition 5 (The class BPP). A decision problem L belongs to the complexity class BPP if there exists a probabilistic Turing machine $\mathcal{M}$ and a constant $0 < \epsilon \leq 1/2$ such that, given an input word w ,

- if $w \in L$, then $\mathbb{P}[\mathcal{M} \text{ accepts } w] \geq 1/2 + \epsilon$; and
- if $w \notin L$, then $\mathbb{P}[\mathcal{M} \text{ accepts } w] \leq 1/2 - \epsilon$.

In the literature, you may see authors taking specific values such as $2/3$ to be the probability of acceptance—as before, this is arbitrary, and in that case such authors simply took $\epsilon = 1/6$ to be their constant. Any positive value can be used in the definition while still retaining the spirit of the class BPP.

Again, drawing a “picture” of the class BPP on a probability line gives us the following, where $0 < \epsilon \leq 1/2$ is our constant:



Note that, now, the probability gap is far more pronounced, and the difference between a correct output and an incorrect output is easier for us to discern depending on our value of ϵ . Indeed, a very nice property of the class BPP is that we can transform an arbitrary BPP-machine into one that has a probability of acceptance extremely close to 1 for all inputs $w \in L$ and extremely close to 0 for all inputs $w \notin L$, while still retaining the overall polynomial runtime of the machine. We will later see the technique that allows us to perform such a transformation.

For now, let us prove a couple more relationships between BPP and our other complexity classes. As before, when we showed that a decision problem in NP is a special case of a decision problem in PP, we can show that a decision problem in P is a special case of a decision problem in BPP.

Theorem 6. $P \subseteq BPP \subseteq PP$.

Proof. We begin by showing that $P \subseteq BPP$. Observe that, for any deterministic Turing machine $\mathcal{M}$ recognizing a language $L \in P$, $\mathcal{M}$ accepts words $w \in L$ with probability $1 = 1/2 + 1/2$ and accepts words $w \notin L$ with probability $0 = 1/2 - 1/2$. Thus, $\mathcal{M}$ is a BPP-machine where $\epsilon = 1/2$.

Showing that $BPP \subseteq PP$ is rather straightforward. From the definition of the class BPP, if $w \in L$, then $\mathbb{P}[\mathcal{M} \text{ accepts } w] \geq 1/2 + \epsilon$ for some $\epsilon > 0$. Therefore, $\mathbb{P}[\mathcal{M} \text{ accepts } w] > 1/2$, which corresponds to the definition of the class PP. We can make a similar argument when $w \notin L$. Thus, every BPP-machine is also a PP-machine. $\square$

Interestingly, while some researchers conjecture that $P = BPP$, nothing is known about the relationship between NP and BPP! However, it seems unlikely that $NP \subseteq BPP$, since this would imply that we can solve NP-complete decision problems reasonably efficiently.

4.3 One-Sided Error: RP

We now consider what happens when we slightly restrict the behaviour of our probabilistic Turing machine so that its computation still runs in polynomial time, but may err in only one of the two outcomes. In this case, our decision problem belongs to the class RP, representing *randomized polynomial time* decision problems.

Definition 7 (The class RP). A decision problem L belongs to the complexity class RP if there exists a probabilistic Turing machine $\mathcal{M}$ and a constant $0 < \epsilon \leq 1/2$ such that, given an input word w ,

- if $w \in L$, then $\mathbb{P}[\mathcal{M} \text{ accepts } w] \geq 1/2 + \epsilon$; and
- if $w \notin L$, then $\mathbb{P}[\mathcal{M} \text{ accepts } w] = 0$.

Just as the definition of the class PP corresponded to decision problems solvable by a Monte Carlo algorithm with *two-sided* error, the definition of the class RP exactly matches the behaviour of a Monte Carlo algorithm with *one-sided* error: in the case where the answer is “yes”, it is incorrect with some bounded probability, and in the case where the answer is “no”, it is always correct.

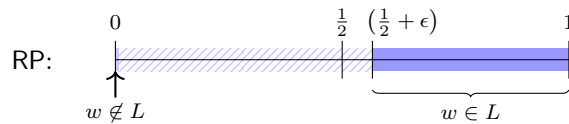
This means that RP-machines are effectively a strengthening of BPP-machines: they behave identically in the case where $w \in L$, but when $w \notin L$ an RP-machine will never accept. As a consequence, we have the following immediate result:

Theorem 8. $\text{RP} \subseteq \text{BPP}$.

Proof. Follows from the definitions of the classes RP and BPP. $\square$

Additionally, as we noted for the class BPP, our choice of value for the probability of acceptance in the definition of RP is arbitrary, and any positive value can serve as the threshold while again retaining the spirit of the class RP.

A “picture” of the class RP on a probability line looks like the following, where $\epsilon > 0$ is some constant:



You may have observed from the definition of the class RP that its acceptance condition for words $w \in L$ is effectively a strengthening of the acceptance condition in our probabilistic definition of the class NP. Therefore, we can draw the following connections between RP and our more familiar complexity classes.

Theorem 9. $\text{P} \subseteq \text{RP} \subseteq \text{NP}$.

Proof. The first relationship showing that $\text{P} \subseteq \text{RP}$ is established in exactly the same way as we established the relationship $\text{P} \subseteq \text{BPP}$.

To see that $\text{RP} \subseteq \text{NP}$, suppose we have a probabilistic Turing machine $\mathcal{M}$ recognizing a language $L \in \text{RP}$. Using the specification of $\mathcal{M}$, we can construct a nondeterministic Turing machine $\mathcal{M}'$ that simply guesses which computation branch to follow instead of using randomization. Since $\mathcal{M}$ accepts words $w \in L$ with probability at least $1/2 + \epsilon$ for some $0 < \epsilon \leq 1/2$, over half of all computation branches lead to an accepting configuration. This means that at least one accepting computation branch exists, and so $\mathcal{M}'$ will accept words $w \in L$ with probability greater than 0. Thus, $L \in \text{NP}$. $\square$

Complementing the Class RP. From our earlier illustration of the class RP, we can clearly see that the acceptance conditions of an RP-machine are not symmetric for words $w \in L$ compared to words $w \notin L$. Ultimately, this means that unlike the classes PP and BPP, we don’t know in general how to construct an RP-machine for a given complement language $\bar{L}$.

We may therefore analogously define a complementary class coRP to handle Monte Carlo algorithms with “other-sided” error in the usual way.

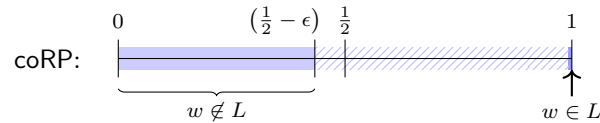
Definition 10 (The class coRP). A decision problem L belongs to the complexity class coRP if there exists a probabilistic Turing machine $\mathcal{M}$ and a constant $0 < \epsilon \leq 1/2$ such that, given an input word w ,

- if $w \in L$, then $\mathbb{P}[\mathcal{M} \text{ rejects } w] = 0$; and
- if $w \notin L$, then $\mathbb{P}[\mathcal{M} \text{ rejects } w] \geq 1/2 + \epsilon$.

Observe that, in our definition of coRP, we have shifted our perspective from probabilities of *acceptance* to probabilities of *rejection*. If some word belongs to the language of a coRP-machine, then it is guaranteed to accept that word, but it may err in rejecting if the word does *not* belong to its language.

Equivalently, we can simply say that a decision problem L belongs to the class coRP if and only if its complementary decision problem $\bar{L}$ belongs to the class RP. Since the acceptance conditions for PP and BPP are symmetric, we have that both $\text{PP} = \text{coPP}$ and $\text{BPP} = \text{coBPP}$. On the other hand, the question of whether RP and coRP are equal remains open.

If we reinterpret the statement “ $\mathbb{P}[\mathcal{M} \text{ rejects } w] \geq 1/2 + \epsilon$ ” to read “ $\mathbb{P}[\mathcal{M} \text{ accepts } w] < 1/2 - \epsilon$ ”, then we can illustrate a “picture” of the class **coRP** to go with that which we had for the class **RP**:



Take special note here that this illustration, like all of the others, depicts probabilities of *acceptance*. A **coRP**-machine will always accept words $w \in L$, while it may (with at most $1/2 - \epsilon$ probability) mistakenly accept words $w \notin L$.

Going along with the results we established in Theorems 8 and 9, it is possible for us to prove both that $\text{coRP} \subseteq \text{BPP}$ and that $\text{P} \subseteq \text{coRP} \subseteq \text{coNP}$. Moreover, if our earlier noted conjecture that $\text{P} = \text{BPP}$ is indeed true, then we would ultimately have a collapse of complexity classes, where $\text{P} = \text{RP} = \text{coRP}$.

4.4 “Zero”-Sided Error: ZPP

For most decision problems solvable by a Monte Carlo algorithm, we need only concern ourselves with one side of the error (that is, erring either on “yes” answers or on “no” answers) because there exists only the one randomized algorithm for the problem. Thus, if we get a “no” answer, for example, we can be confident in its correctness, while we may need to take extra steps to verify a “yes” answer.

For a certain subset of decision problems, however, we can devise *two* Monte Carlo algorithms to correctly handle “yes” answers and “no” answers, respectively. In this case, since each algorithm errs in a different side of the output, combining these algorithms and ignoring the erroneous side of the output gives us a way of always obtaining the correct answer.

If a decision problem has two randomized algorithms handling either side of the output in this way, then we say it belongs to the class **ZPP** of *zero-error probabilistic polynomial time* decision problems.

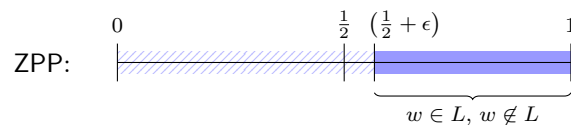
Definition 11 (The class **ZPP**). A decision problem L belongs to the complexity class **ZPP** if there exists a probabilistic Turing machine $\mathcal{M}$ and a constant $0 < \epsilon \leq 1/2$ such that, given an input word w ,

- if $w \in L$, then $\mathbb{P}[\mathcal{M} \text{ accepts } w] \geq 1/2 + \epsilon$, and otherwise halts in a “don’t know” state; and
- if $w \notin L$, then $\mathbb{P}[\mathcal{M} \text{ rejects } w] \geq 1/2 + \epsilon$, and otherwise halts in a “don’t know” state.

In our definition, the purpose of the “don’t know” answer is to allow us to ignore the erroneous side of each algorithm’s output, as we noted earlier.

Since a **ZPP**-machine always produces either the correct answer or a “don’t know” answer, we can be confident that the **ZPP**-machine is correct if it either accepts or rejects its input word. However, since it is possible to receive “don’t know” as the output, we can’t say with confidence how long it will take for a **ZPP**-machine to produce its accept-or-reject answer. The class **ZPP** therefore corresponds to the set of decision problems solvable by a Las Vegas algorithm.

Once more, we can draw a “picture” of the class **ZPP** on a probability line, where $\epsilon > 0$ is some constant. However, here we simply interpret the probability line as illustrating the probability of getting the correct answer (i.e., if $w \in L$, we consider the probability of accepting, while if $w \notin L$, we consider the probability of rejecting):



Since the probabilities of a **ZPP**-machine accepting and rejecting its input word are symmetric, much like we had with the classes **PP** and **BPP**, it is the case that $\text{ZPP} = \text{coZPP}$.

Directly from the definition, we see that the class ZPP is effectively the set of all decision problems that can be solved by both an RP -machine and a coRP -machine. Indeed, this is what we mean by the decision problem having two randomized algorithms handling either side of the output. We can prove this fact formally as follows.

Theorem 12. $\text{ZPP} = \text{RP} \cap \text{coRP}$.

Proof. To show that $\text{ZPP} \subseteq \text{RP} \cap \text{coRP}$, let $\mathcal{M}_Z$ be a ZPP -machine recognizing a language L . We can construct an RP -machine $\mathcal{M}_R$ recognizing the same language L by changing all “don’t know” outputs of $\mathcal{M}_Z$ to “reject” outputs. Now, $\mathcal{M}_R$ accepts with probability $1/2 + \epsilon$ if $w \in L$ and rejects with probability 1 otherwise, so $\text{ZPP} \subseteq \text{RP}$. Similarly, we can construct a coRP -machine $\mathcal{M}_{\bar{R}}$ recognizing L by changing all “reject” outputs of $\mathcal{M}_Z$ to “accept” outputs and changing all “accept” and “don’t know” outputs of $\mathcal{M}_Z$ to “reject” outputs. Thus, $\text{ZPP} \subseteq \text{coRP}$ as well.

In the other direction, to show that $\text{RP} \cap \text{coRP} \subseteq \text{ZPP}$, let $\mathcal{M}_R$ be an RP -machine for some language L and let $\mathcal{M}_{\bar{R}}$ be an RP -machine for the complement language $\bar{L}$. We can construct a ZPP -machine $\mathcal{M}_Z$ for L in the following way by simulating the computations of both $\mathcal{M}_R$ and $\mathcal{M}_{\bar{R}}$ and taking the result to be (z_1, z_2) , where $z_i \in \{\text{accept, reject}\}$ for $i \in \{1, 2\}$. The machine then makes its decision based on which result is produced:

- If the result is (accept, reject), then $\mathcal{M}_Z$ accepts.
- If the result is (reject, accept), then $\mathcal{M}_Z$ rejects.
- If the result is (reject, reject), then $\mathcal{M}_Z$ outputs “don’t know”.
- The result (accept, accept) is impossible.

In the first case, $\mathcal{M}_Z$ accepts with probability at least $1/2 + \epsilon$, and in the second case, $\mathcal{M}_Z$ rejects with probability at least $1/2 + \epsilon$. Thus, $\text{RP} \cap \text{coRP} \subseteq \text{ZPP}$.

Since both directions of the subset inclusion hold, we have that $\text{ZPP} = \text{RP} \cap \text{coRP}$ as desired. $\square$

As a nice consequence, the proof of Theorem 12 gives us a method of converting two Monte Carlo algorithms to one Las Vegas algorithm, which nicely complements our Las-Vegas-to-Monte-Carlo conversion procedure given by Theorem 2.

We can of course establish a number of straightforward relationships between the class ZPP and other complexity classes. For instance, as a consequence of Theorem 9, we have that $\text{ZPP} \subseteq \text{NP} \cap \text{coNP}$. In addition to this, since we know that both $\text{P} \subseteq \text{RP}$ and $\text{P} \subseteq \text{coRP}$, we have that $\text{P} \subseteq \text{RP} \cap \text{coRP}$ and so, with Theorem 12, we can conclude that $\text{P} \subseteq \text{ZPP}$.

5 The Randomized Complexity Hierarchy

Combining all of our relationships between each of the randomized complexity classes we defined here, and including relationships between our fundamental complexity classes such as P , NP , and PSPACE , we can obtain a finer-grained randomized complexity hierarchy that refines our fundamental complexity hierarchy by “filling in” the relationships between P and PSPACE . This randomized complexity hierarchy is depicted in Figure 2.

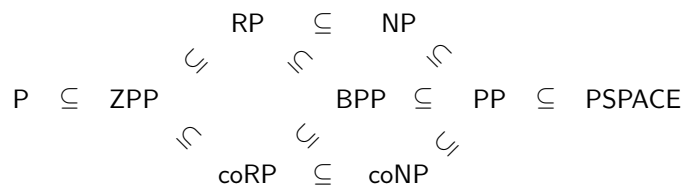


Figure 2: The randomized complexity hierarchy.