

*The art of art, the glory of expression
and the sunshine of the light of letters is simplicity.
Nothing is better than simplicity . . .
nothing can make up for excess or for the lack of definiteness.*

— Walt Whitman, Preface to *Leaves of Grass* (1855)

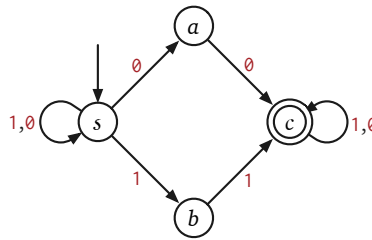
*Freedom of choice
Is what you got.
Freedom from choice
Is what you want.*

— Devo, “Freedom of Choice”, *Freedom of Choice* (1980)

4 Nondeterminism

4.1 Nondeterministic State Machines

The following diagram shows something that looks like a finite-state machine over the alphabet $\{0, 1\}$, but on closer inspection, it is not consistent with our earlier definitions. On one hand, there are two transitions out of s for each input symbol. On the other hand, states a and b are each missing an outgoing transition.



A nondeterministic finite-state automaton

This object is an example of a **nondeterministic finite-state automaton**, or **NFA**, so named because its behavior is not uniquely *determined* by the input string. Formally, every NFA has five components:

- An arbitrary finite set Σ , called the **input alphabet**.
- Another arbitrary finite set Q , whose elements are called **states**.
- An arbitrary **transition** function $\delta : Q \times \Sigma \rightarrow 2^Q$.
- A **start state** $s \in Q$.
- A subset $A \subseteq Q$ of **accepting states**.

The only difference from the formal definition of *deterministic* finite-state automata is the output type of the transition function. In a DFA, the transition function always returns a single state; in an NFA, the transition function returns always a *set* of states, which could be empty, or all of Q , or anything in between.¹

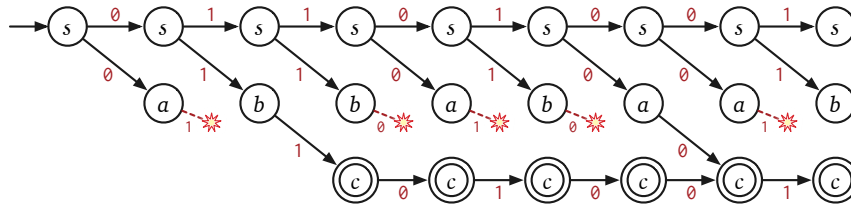
¹“ 2^Q ” is mnemonic notation for the *power set* of Q , which is the set of all subsets of Q .

Just like DFAs, the behavior of an NFA is governed by an *input string* $w \in \Sigma^*$, which the machine reads one symbol at a time, from left to right. Unlike DFAs, however, an NFA does not maintain a single current state, but rather a *set* of current states. Whenever the NFA reads a symbol a , its set of current states changes from C to $\delta(C, a) := \bigcup_{q \in C} \delta(q, a)$. After all symbols have been read, the NFA **accepts** w if its current state set contains *at least one* accepting state and **rejects** w otherwise. In particular, if the set of current states ever becomes empty, it will stay empty forever, and the NFA will reject.

For example, given the input string **01101001**, our example NFA's set of current states evolves as follows.

$$\{s\} \xrightarrow{0} \{s, a\} \xrightarrow{1} \{s, b\} \xrightarrow{1} \{s, b, c\} \xrightarrow{0} \{s, a, c\} \xrightarrow{1} \{s, b, c\} \xrightarrow{0} \{s, a, c\} \xrightarrow{0} \{s, a, c\} \xrightarrow{1} \{s, b, c\}$$

This sequence of transitions is drawn more explicitly in the figure below. (Notice the explosions whenever the NFA reads a **1** in state a , or reads a **0** in state b .) Because the final set of states $\{a, b, c\}$ contains the accept state c , our NFA accepts this input string.



Running our example NFA on the input string **01101001**.

More formally, for any NFA $(\Sigma, Q, s, A, \delta)$, we recursively define an *extended transition function* $\delta^*: Q \times \Sigma^* \rightarrow 2^Q$ that takes a state and a *string* as input, and returns a set of states as output:

$$\delta^*(q, w) := \begin{cases} \{q\} & \text{if } w = \varepsilon, \\ \bigcup_{r \in \delta(q, a)} \delta^*(r, x) & \text{if } w = ax. \end{cases}$$

For example, in our example NFA, we have $\delta^*(s, \mathbf{01101001}) = \{s, b, c\}$ and $\delta^*(s, \mathbf{01010}) = \{s, a\}$. Finally, the NFA **accepts** a string $w \in \Sigma^*$ if and only if $\delta^*(s, w) \cap A \neq \emptyset$.

We sometimes abuse notation by defining $\delta(P, a) = \bigcup_{q \in P} \delta(q, a)$ and $\delta^*(P, w) = \bigcup_{q \in P} \delta^*(q, w)$ for any subset of states $P \subseteq Q$. With this new notation, we have

$$\delta^*(q, w) := \begin{cases} \{q\} & \text{if } w = \varepsilon, \\ \delta^*(\delta(q, a), x) & \text{if } w = ax. \end{cases}$$

We can equivalently define an NFA as *any* directed graph whose vertices are the states Q , and whose edges are labeled with symbols from Σ . Unlike DFAs, we do not require that every vertex has exactly one outgoing edge with each label; it may have several such edges or none. An NFA accepts a string w if the graph contains *at least one* walk from the start state to an accepting state whose label is the string w .

4.2 Intuition

There are at least three useful ways to think about non-determinism.

Clairvoyance. Whenever an NFA reads symbol a in state q , it *chooses* the next state from the set $\delta(q, a)$, always *magically* choosing a state that leads to the NFA accepting the input string, unless no such choice is possible. Even if there is only one sequence of choices out of 14,000,605 that leads the machine to an accepting state, the NFA magically makes that choice.² As the BSD fortune file put it, “Nondeterminism means never having to say you are wrong.”³ Of course real machines can’t actually look into the future; that’s why I used the word “magic”.

Parallel threads. An arguably more “realistic” view is that when an NFA reads symbol a in state q , it spawns an independent execution thread for each state in $\delta(q, a)$. In particular, if $\delta(q, a)$ is empty, the current thread simply dies. The NFA accepts if *at least one* thread is in an accepting state after it reads the last input symbol. This is the intuition illustrated in the figure on the previous page.

Equivalently, we can imagine that when an NFA reads symbol a in state q , it branches into several parallel universes, one for each state in $\delta(q, a)$. If $\delta(q, a)$ is empty, the NFA destroys the universe (including itself) to protect the Sacred Timeline. Similarly, if the NFA finds itself in a non-accepting state when the input ends, the NFA destroys the universe to protect the Sacred Timeline. Thus, when the input is gone, only universes in which the NFA somehow chose a path to an accept state still exist. One slight disadvantage of this metaphor is that if an NFA reads a string that is not in its language, it destroys *all* universes.

Proofs/oracles. Finally, we can treat NFAs not as a mechanism for *computing* something, but as a mechanism for *verifying proofs*. If we want to *prove* that a string w contains one of the suffixes **00** or **11**, it suffices to demonstrate a single walk in our example NFA that starts at s and ends at c , and whose edges are labeled with the symbols in w . Equivalently, whenever the NFA faces a nontrivial choice, the prover can simply tell the NFA which state to move to next.

This intuition can be formalized as follows. Consider a *deterministic* finite state machine whose input alphabet is the product $\Sigma \times \Omega$ of an *input* alphabet Σ and an *oracle* alphabet Ω . Equivalently, we can imagine that this DFA reads simultaneously from two strings of the same length: the *input* string w and the *oracle* string ω . In either formulation, the transition function has the form $\delta: Q \times (\Sigma \times \Omega) \rightarrow Q$. As usual, this DFA accepts the pair $(w, \omega) \in (\Sigma \times \Omega)^*$ if and only if $\delta^*(s, (w, \omega)) \in A$. Finally, M **nondeterministically accepts** the string $w \in \Sigma^*$ if there is an oracle string $\omega \in \Omega^*$ such that $|\omega| = |w|$ and $(w, \omega) \in L(M)$.

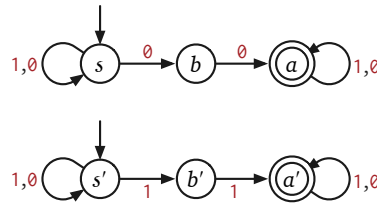
4.3 Multiple Start States

It is arguably more natural (but a bit non-standard) to allow NFAs to have a *set* S of start states instead of a single start state s . Such an NFA would non-deterministically choose its actual start state before reading any input symbols. Thus, an NFA with multiple start states accepts a string w if and only if there is a sequence of transitions consistent with w from *some* start state to *some* accepting state, or more formally if $\delta^*(S, w) \cap A \neq \emptyset$.

For example, consider the following NFA (yes, just one) with multiple start states, each indicated with its own incoming arrow.

²Yes, even if it means Black Widow and Iron Man have to die.

³This sentence is a riff on a horrible aphorism that was (sadly) popular in the US in the 70s and 80s. Fortunately, everyone seems to have forgotten the original saying, except maybe for that one time it was mocked on *The Simpsons*. Ah, who am I kidding? Nobody remembers *The Simpsons* either.



An NFA with two start states.

This NFA accepts the string 01101 by nondeterministically choosing the start state s' and then nondeterministically choose the following sequence of transitions:

$$s' \xrightarrow{0} s' \xrightarrow{1} b' \xrightarrow{1} a' \xrightarrow{0} a' \xrightarrow{1} a'$$

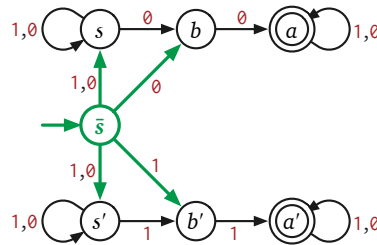
Alternatively, we can track the current set of states as follows:

$$\{s, s'\} \xrightarrow{0} \{s, b, s'\} \xrightarrow{1} \{s, s', b'\} \xrightarrow{1} \{s, s', b', a'\} \xrightarrow{0} \{s, b, s', a'\} \xrightarrow{1} \{s, s', b', a'\}$$

The NFA accepts the string 01101 because the final set of states contains the accepting state a' .

Given any NFA with multiple start states S , we can construct an equivalent NFA with only one start state $\bar{s}$ as follows:

- Add a new state $\bar{s}$.
- Copy all outgoing transitions from every start state in S to $\bar{s}$.
- Declare $\bar{s}$ to be an accepting state if and only if S contains an accepting state.
- Declare $\bar{s}$ to be the only start state.



An equivalent NFA with only one start state.

More formally, given any multiple-start NFA $M = (\Sigma, Q, S, A, \delta)$, we can construct an equivalent single-start NFA $\bar{M} = (\Sigma, \bar{Q}, \bar{s}, \bar{A}, \bar{\delta})$ by defining

$$\bar{Q} = Q \cup \{\bar{s}\}$$

$$\bar{A} = \begin{cases} A & \text{if } S \cap A = \emptyset \\ A \cup \{\bar{s}\} & \text{otherwise} \end{cases}$$

$$\bar{\delta}(\bar{s}, a) = \bigcup_{s \in S} \delta(s, a) \quad \text{for all } a \in \Sigma$$

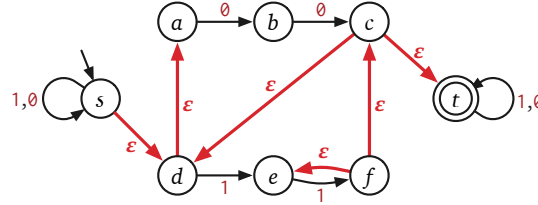
$$\bar{\delta}(q, a) = \delta(q, a) \quad \text{for all } q \in Q \text{ and } a \in \Sigma$$

Straightforward definition-chasing now implies that M and $\bar{M}$ accept exactly the same language. Thus, whenever we reason about or design NFAs, we are free to either allow or forbid multiple start states, whichever is more convenient for the task at hand.

4.4 ϵ -Transitions

It is also fairly common for NFAs to include so-called **ϵ -transitions**, which allow the machine to change state without reading an input symbol. An NFA with ϵ -transitions accepts a string w if and only if there is a sequence of transitions $s \xrightarrow{a_1} q_1 \xrightarrow{a_2} q_2 \xrightarrow{a_3} \dots \xrightarrow{a_\ell} q_\ell$ where the final state q_ℓ is accepting, each a_i is either ϵ or a symbol in Σ , and $a_1 a_2 \dots a_\ell = w$.

For example, consider the following NFA with ϵ -transitions. (For this example, we indicate the ϵ -transitions using large red arrows; we won't normally do that.) This NFA deliberately has more ϵ -transitions than necessary.



A (rather silly) NFA with ϵ -transitions

This NFA has a single start state s . If the input string is **100111**, the machine might non-deterministically choose the following transitions and then accept.

$$s \xrightarrow{\epsilon} s \xrightarrow{1} s \xrightarrow{\epsilon} d \xrightarrow{\epsilon} a \xrightarrow{0} b \xrightarrow{0} c \xrightarrow{\epsilon} d \xrightarrow{1} e \xrightarrow{1} f \xrightarrow{\epsilon} e \xrightarrow{1} f \xrightarrow{\epsilon} c \xrightarrow{\epsilon} t$$

More formally, the transition function in an NFA with ϵ -transitions has a slightly larger domain $\delta: Q \times (\Sigma \cup \{\epsilon\}) \rightarrow 2^Q$. The **ϵ -reach** of a state $q \in Q$ consists of all states r that satisfy one of the following recursive conditions:

- either $r = q$,
- or $r \in \delta(q', \epsilon)$ for some state q' in the ϵ -reach of q .

In other words, r is in the ϵ -reach of q if there is a (possibly empty) sequence of ϵ -transitions leading from q to r . For instance, in the example NFA above, the ϵ -reach of state f is the set $\{a, c, d, e, f, g\}$, but the ϵ -reach of state a is just $\{a\}$.

Now we redefine the extended transition function $\delta^*: Q \times \Sigma^* \rightarrow 2^Q$, which transitions on arbitrary strings, as follows:

$$\delta^*(p, w) := \begin{cases} \epsilon\text{-reach}(p) & \text{if } w = \epsilon, \\ \bigcup_{r \in \epsilon\text{-reach}(p)} \bigcup_{q \in \delta(r, a)} \delta^*(q, x) & \text{if } w = ax. \end{cases}$$

If we extend our earlier notational abuse by writing $\delta(P, a) = \bigcup_{q \in P} \delta(q, a)$ and $\delta^*(P, w) = \bigcup_{q \in P} \delta^*(q, w)$ and $\epsilon\text{-reach}(P) = \bigcup_{q \in P} \epsilon\text{-reach}(q)$ for any subset of states $P \subseteq Q$, this definition simplifies as follows:

$$\delta^*(p, w) := \begin{cases} \epsilon\text{-reach}(p) & \text{if } w = \epsilon, \\ \delta^*(\delta(\epsilon\text{-reach}(p), a), x) & \text{if } w = ax. \end{cases}$$

Intuitively, the modified δ^* function tries *all* possible ways of interspersing ε -transitions before, after, and between regular symbol transitions. For instance, for our example NFA at the top of this page, we can compute $\delta^*(s, 11) = \{s, a, c, d, e, f, g\}$ as follows:

$$s \xrightarrow{\varepsilon\text{-reach}} \{s, a, d\} \xrightarrow{1} \{s, e\} \xrightarrow{\varepsilon\text{-reach}} \{s, a, d, e\} \xrightarrow{1} \{s, e, f\} \xrightarrow{\varepsilon\text{-reach}} \{s, a, c, d, e, f, g\}.$$

Finally, as usual, an NFA with ε -transitions accepts a string w if and only if $\delta^*(s, w)$ contains at least one accepting state. Thus, our example NFA accepts the string 11 , because $\delta^*(s, 11) = \{s, a, c, d, e, f, g\}$ contains the accepting state g . An NFA with ε -transitions *and multiple start states* accepts w if and only if $\delta^*(S, w) \cap A \neq \emptyset$.

4.5 Removing ε -Transitions

Although it may appear at first that ε -transitions give us a more powerful set of machines, NFAs with and without ε -transitions are actually equivalent. Given an NFA $M = (\Sigma, Q, s, A, \delta)$ with ε -transitions, we can construct an equivalent without ε -transitions but with multiple start states as follows:

- First, declare every state in the ε -reach of s to be a start state. This modification does not change the accepted language.
- For every transition pair $p \xrightarrow{a} q \xrightarrow{\varepsilon} r$, add a direct transition $p \xrightarrow{a} r$. This addition does not change the accepted language.
- When no more of the previous modifications are possible, delete all ε -transitions. This modification does not change the accepted language.

More formally, the new NFA $M' = (\Sigma, Q', S', A', \delta')$ is defined as follows:

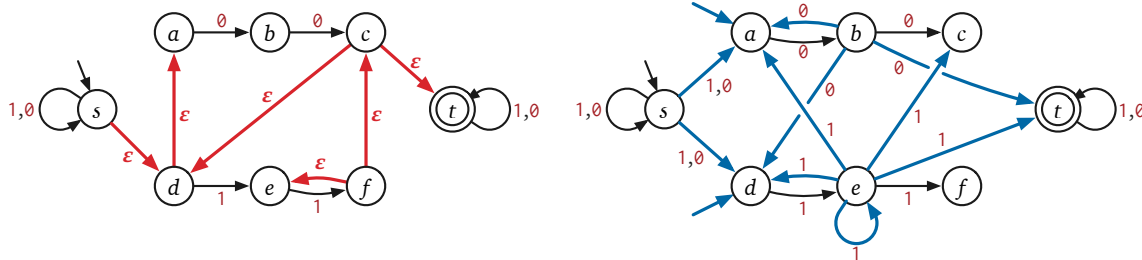
$$\begin{aligned} Q' &:= Q \\ S' &:= \varepsilon\text{-reach}(s) \\ A' &:= A \\ \delta'(q, a) &:= \varepsilon\text{-reach}(\delta(q, a)) \end{aligned}$$

(If necessary, we can further modify M' to have a single start state.) Tedious inductive definition-chasing now implies that M and M' accept exactly the same language. Thus, whenever we reason about or design NFAs, we are free to either allow or forbid ε -transitions, whichever is more convenient for the task at hand.

For example, our previous NFA with ε -transitions can be transformed into an equivalent NFA without ε -transitions, as shown in the figure on the next page. The transitions added by the transformation are drawn in blue. For example, the chain of transitions

$$e \xrightarrow{1} f \xrightarrow{\varepsilon} c \xrightarrow{\varepsilon} d \xrightarrow{\varepsilon} a$$

in the NFA on the left becomes a single transition $e \xrightarrow{1} a$ in the NFA on the right. The output NFA on the right has three start states, as well as two states c and f that are clearly redundant, but whatever.



A (rather silly) NFA with ϵ -transitions, and an equivalent (still silly) NFA without ϵ -transitions.

4.6 Kleene's Theorem

We are now finally in a position to prove the following fundamental theorem, first observed by Stephen Kleene in 1951:

Theorem 4.1. *A language L can be described by a regular expression if and only if L is the language accepted by a DFA.*

We will prove Kleene's fundamental theorem in four stages:

- Every DFA can be transformed into an equivalent NFA.
- Every NFA can be transformed into an equivalent DFA.
- Every regular expression can be transformed into an equivalent NFA.
- Every NFA can be transformed into an equivalent regular expression.

The first of these four transformations is completely trivial; a DFA is just a special type of NFA where the transition function always returns a single state. Unfortunately, the other three transformations require a bit more work.

4.7 NFA to DFA: The Subset Construction

In the parallel-thread model of NFA execution, an NFA does not have a single current state, but rather a *set* of current states. The evolution of this set of states is **determined** by a modified transition function $\delta': 2^Q \times \Sigma \rightarrow 2^Q$, defined by setting $\delta'(P, a) := \bigcup_{p \in P} \delta(p, a)$ for any set of states $P \subseteq Q$ and any symbol $a \in \Sigma$. When the NFA finishes reading its input string, the NFA accepts if and only if the current set of NFA states intersects the set A of NFA accepting states.

This formulation actually makes the NFA completely **deterministic**! We have just shown that any NFA $M = (\Sigma, Q, S, A, \delta)$ with multiple start states, but without ϵ -transitions, is equivalent to the DFA $M' = (\Sigma, Q', s', A', \delta')$ defined as follows:

$$\begin{aligned}
 Q' &:= 2^Q \\
 s' &:= S \\
 A' &:= \{q' \subseteq Q \mid q' \cap A \neq \emptyset\} \\
 \delta'(q', a) &:= \bigcup_{p \in q'} \delta(p, a) \quad \text{for all } q' \subseteq Q \text{ and } a \in \Sigma.
 \end{aligned}$$

Similarly, any NFA with ε -transitions is equivalent to a DFA defined as follows:

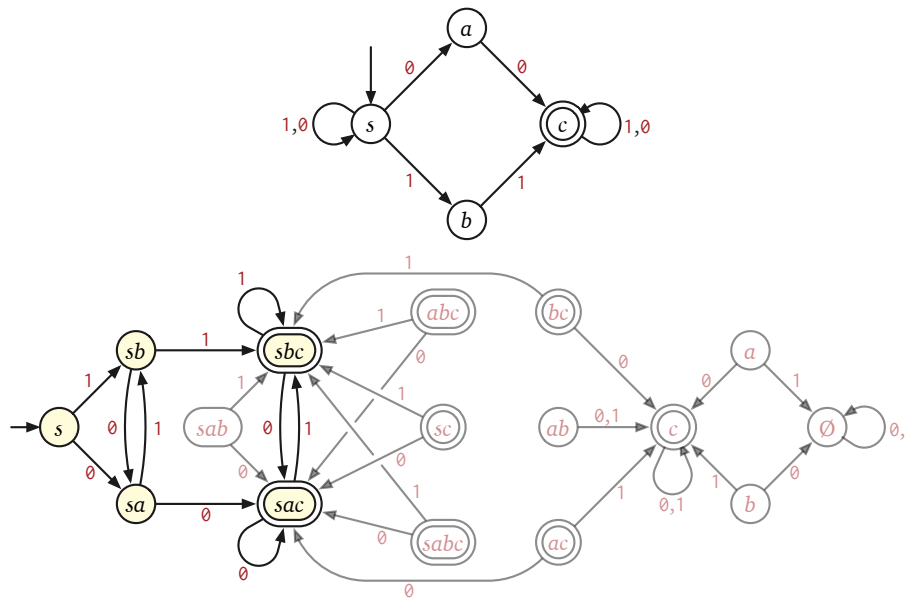
$$\begin{aligned}
 Q' &:= 2^Q \\
 s' &:= \varepsilon\text{-reach}(S) \\
 A' &:= \{q' \subseteq Q \mid q' \cap A \neq \emptyset\} \\
 \delta'(q', a) &:= \bigcup_{p \in q'} \bigcup_{r \in \delta(p, a)} \varepsilon\text{-reach}(r) \quad \text{for all } q' \subseteq Q \text{ and } a \in \Sigma.
 \end{aligned}$$

Or more concisely:

$$\delta'(q', a) := \varepsilon\text{-reach}(\delta(q', a)) \quad \text{for all } q' \subseteq Q \text{ and } a \in \Sigma.$$

This translation from NFA to equivalent DFA is usually called the **subset construction** (or the **power-set construction**), but that name is somewhat misleading. It's not a "construction" so much as a change in perspective.

For example, the 4-state NFA on the first page of this note is equivalent to the following 16-state DFA. To simplify notation, I've named each DFA state using a simple string, omitting the braces and commas from the corresponding subset of NFA states; for example, DFA state *sbc* corresponds to the subset $\{s, b, c\}$ of NFA states.



A 4-state NFA and an equivalent 16-state DFA defined by the subset construction.
States are unreachable from the start state *s* are grayed out.

An obvious disadvantage of this "construction" is that it (usually) leads to DFAs with far more states than necessary, in part because many states cannot even be reached from the start state. In the example above, only five of the sixteen states of the DFA are reachable from *s*.

Caveat lector: Talking and writing about both an NFA and its equivalent DFA at the same time can be confusing. In these notes I carefully add "DFA" or "NFA" in front of the name of any component of either machine, writing "NFA state" or "DFA state" instead of just "state", for example. I also consistently use variables without primes to denote components of the

NFA, and variables with primes to denote correspond components of the equivalent DFA. For example, S denotes the set of NFA start states, and s' denotes the equivalent single DFA start state. Similarly, q denotes a single NFA state, while q' denotes a single DFA state (or equivalently(!) a set of NFA states).

Incremental Subset Construction

Instead of building the entire powerset DFA and then discarding unreachable states, we can avoid the unreachable states from the beginning by constructing a DFA incrementally, essentially by performing a breadth-first search of the implicitly-represented powerset DFA graph.

To execute this algorithm by hand, we prepare a table with $2|\Sigma| + 2$ columns, with one row for each DFA state we discover. In order, these columns record the following information:

- The DFA state (as a subset of NFA states)
- Whether the DFA state is accepting (that is, whether the ε -reach intersects A)
- The output of the transition function δ for each symbol in Σ .
- The ε -reaches of the corresponding subsets of NFA states

We start the algorithm by writing the ε -reach of the NFA start state(s) in the first row and first column. Whenever we discover an unexplored DFA state in one of the $|\Sigma|$ ε -reach columns, we copy it to the left column in a new row. To reduce notational clutter, we write all subsets of NFA states without braces or commas.

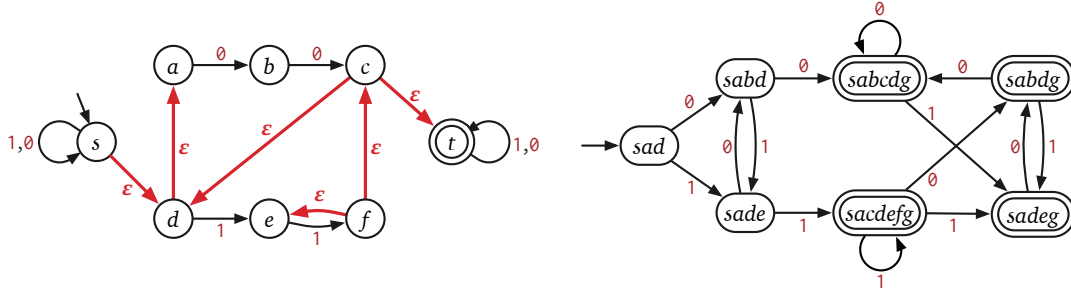
For example, given our rather silly NFA with ε -transitions from Section 4.4, a naïve subset construction would produce a DFA with 256 states,⁴ but the *incremental* subset construction produces a DFA with only *seven* states, described by the following table and illustrated on the next page:

q'	$q' \in A'?$	$\delta'(q', \emptyset)$	$\varepsilon\text{-reach}(\delta'(q', \emptyset))$	$\delta'(q', 1)$	$\varepsilon\text{-reach}(\delta'(q', 1))$
$\varepsilon\text{-reach}(a) = sad$		sb	$sabd$	se	$sade$
$sabd$		sbc	$sabcdg$	se	$sade$
$sade$		sb	$sabd$	sef	$sacdefg$
$sabcdg$	✓	$sbcg$	$sabcdg$	seg	$sadeg$
$sacdefg$	✓	$sbfg$	$sabdg$	$sefg$	$sacdefg$
$sadeg$	✓	$sbfg$	$sabdg$	$sefg$	$sacdefg$
$sabdg$	✓	$sbcg$	$sabcdg$	seg	$sadeg$

The incremental construction algorithm fills in the top row of the table as follows:

- We have $\varepsilon\text{-reach}(s) = \{s, a, d\}$, so we write sad in the first column.
- None of the NFA states $\{s, a, d\}$ is an accepting state, so $\{s, a, d\}$ is not an accepting state of the DFA, so we do *not* check the second column.
- Next, $\delta'(\{s, a, d\}, \emptyset) = \delta(s, \emptyset) \cup \delta(a, \emptyset) \cup \delta(d, \emptyset) = \{s\} \cup \{b\} \cup \emptyset = \{s, b\}$, so we write sb in the third column.

⁴In fact, we can immediately quickly discard any subset of NFA states that is not closed under ε -transitions, but that still leaves a *lot* of DFA states.



An eight-state NFA with ϵ -transitions, and the output of the incremental subset construction for that NFA.

- Next, $\epsilon\text{-reach}(\{s, b\}) = \epsilon\text{-reach}(s) \cup \epsilon\text{-reach}(b) = \{s, a, d\} \cup \{b\} = \{s, a, b, d\}$, so we write *sabd* in the fourth column. Because *sabd* does not already appear in the first column, we have discovered a new DFA state! We start a new row for DFA state *sabd*.
- Finally, $\delta'(\{s, a, d\}, 1) = \delta(s, 1) \cup \delta(a, 1) \cup \delta(d, 1) = \{s\} \cup \emptyset \cup \{e\} = \{s, e\}$, so we write *se* in the fourth column, and we start a new row for the new DFA state *se*.

We now have two new rows to fill in, corresponding to DFA-states *sabd* and *sade*. The algorithm continues filling in rows (and discovering new rows) until all rows are filled.

Although it avoids *unreachable* states, the incremental subset algorithm can still produce DFAs with more states than necessary. For this example NFA, for example, the resulting DFA keeps looking for **00** and **11** substrings even after it's already found one. When the NFA finds both **00** and **11** as substrings, it doesn't kill all the other parallel execution threads, *because NFAs can't do that*. NFAs often have significantly fewer states than equivalent DFAs, but that efficiency also makes them kind of stupid.

Alternative “ ϵ -first” Construction

In an earlier version of this note, my definition of the subset construction reversed the order of chasing real transitions and chasing ϵ -transitions. In particular, instead of chasing ϵ -transitions before the DFA starts, this alternative formulation chases ϵ -transitions in the acceptance criterion when the input ends.

$$Q' := 2^Q$$

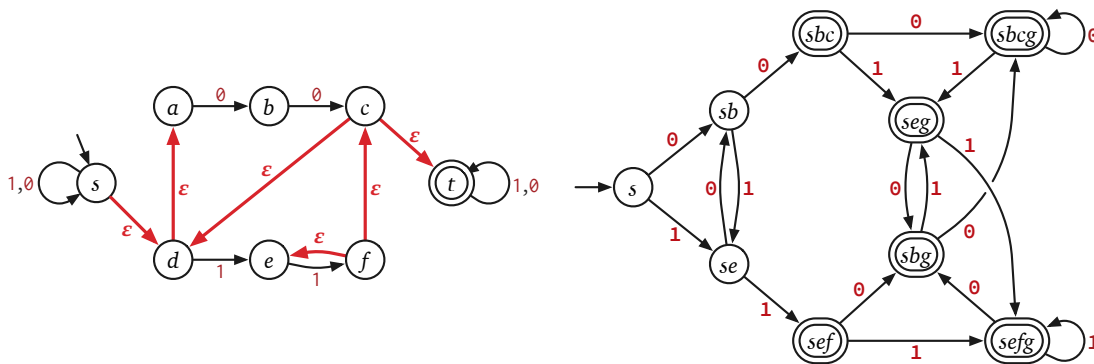
$$s' := s$$

$$A' := \{q' \subseteq Q \mid \epsilon\text{-reach}(q') \cap A \neq \emptyset\}$$

$$\delta'(q', a) := \delta(\epsilon\text{-reach}(q'), a) \quad \text{for all } q' \subseteq Q \text{ and } a \in \Sigma.$$

This “ ϵ -first” subset construction also yields an alternative incremental subset construction algorithm. For each row, we compute the ϵ -reach of the set q' NFA states named in that row, check whether $\epsilon\text{-reach}(q')$ contains an accepting state, evaluates the transition function on $\epsilon\text{-reach}(q')$, and then start new rows for any new DFA states we discover. This variant of the incremental subset construction tends to produce DFAs with more states, but shorter state names.

q'	$\varepsilon\text{-reach}(q')$	$q' \in A'?$	$\delta(\varepsilon\text{-reach}(q'), 0)$	$\delta(\varepsilon\text{-reach}(q'), 1)$
s	sad		sb	se
sb	$sabd$		sbc	se
se	$sade$		sb	sef
sbc	$sabcdg$	✓	$sbcg$	seg
sef	$sacdefg$	✓	sb	$sefg$
$sbcg$	$sabcdg$	✓	$sbcg$	seg
seg	$sadeg$	✓	sb	$sefg$
sb	$sabd$	✓	$sbcg$	seg
$sefg$	$sacdefg$	✓	sb	$sefg$



An eight-state NFA with ε -transitions, and the output of the “ ε -first” incremental subset construction for that NFA.

4.8 Regular Expression to NFA: Thompson’s Algorithm

We now turn to the core of Kleene’s theorem, which claims that regular languages (described by regular expressions) and automatic languages (accepted by finite-state automata) are the same.

Lemma 4.2. *Every regular language is accepted by a nondeterministic finite-state automaton.*

Proof: The following construction was first described by Ken Thompson in 1968. Thompson’s algorithm actually proves a stronger statement: For any regular language L , there is an NFA that accepts L that has a single start state s and a single accepting state t , where these two states are distinct.

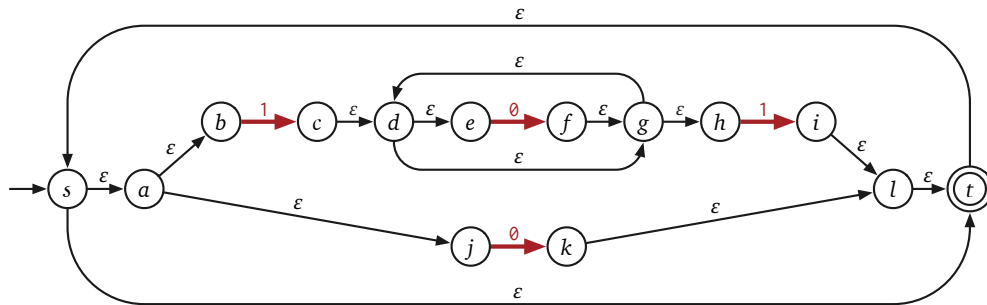
Let R be an arbitrary regular expression over an arbitrary finite alphabet Σ . Assume that for any sub-expression S of R , the language described by S is accepted by an NFA with a unique accepting state distinct from its unique start state, which we denote pictorially by $\circ \text{ } s \text{ } \odot$. There are six cases to consider—three base cases and three recursive cases—mirroring the recursive definition of a regular expression.

- If $R = \emptyset$, then $L(R) = \emptyset$ is accepted by the trivial NFA $\circ \text{ } \odot$.
- If $R = \varepsilon$, then $L(R) = \{\varepsilon\}$ is accepted by a different trivial NFA $\circ \xrightarrow{\varepsilon} \odot$.

- If $R = w$ for some non-empty string $w \in \Sigma^+$, then $L(R) = \{w\}$ is accepted by an NFA consisting of a single chain of transitions from s to t , labeled by the symbols in w in order. For example, the language $\{101\}$ is accepted by the NFA
- Suppose $R = ST$ for some regular expressions S and T . The inductive hypothesis implies that the languages $L(S)$ and $L(T)$ are accepted by NFAs and , respectively. Then $L(R) = L(ST) = L(S) \cdot L(T)$ is accepted by the NFA , built by connecting the two component NFAs in series with a single ϵ -transition.
- Suppose $R = S + T$ for some regular expressions S and T . The inductive hypothesis implies that the languages $L(S)$ and $L(T)$ are accepted by NFAs and , respectively. Then $L(R) = L(S + T) = L(S) \cup L(T)$ is accepted by the NFA , built by connecting the two component NFAs in parallel to new start and accept states using ϵ -transitions.
- Finally, suppose $R = S^*$ for some regular expression S . The inductive hypothesis implies that the language $L(S)$ is accepted by an NFA . Then the language $L(R) = L(S^*) = L(S)^*$ is accepted by the NFA

In all cases, the language $L(R)$ is accepted by an NFA with a unique accepting state, which is different from its unique start state, as claimed. $\square$

As an example, given the regular expression $(0 + 10^*1)^*$ of strings containing an even number of 1s, Thompson's algorithm produces a 14-state NFA shown below. As this example shows, Thompson's algorithm tends to produce NFAs with many redundant states. Fortunately, just as there are for DFAs, there are algorithms that can reduce any NFA to an equivalent NFA with the smallest possible number of states.⁵

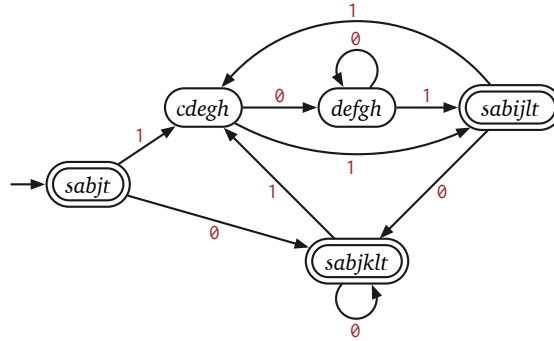


The NFA constructed by Thompson's algorithm for the regular expression $(0 + 10^*1)^*$.
The four non- ϵ -transitions are drawn with bold red arrows for emphasis.

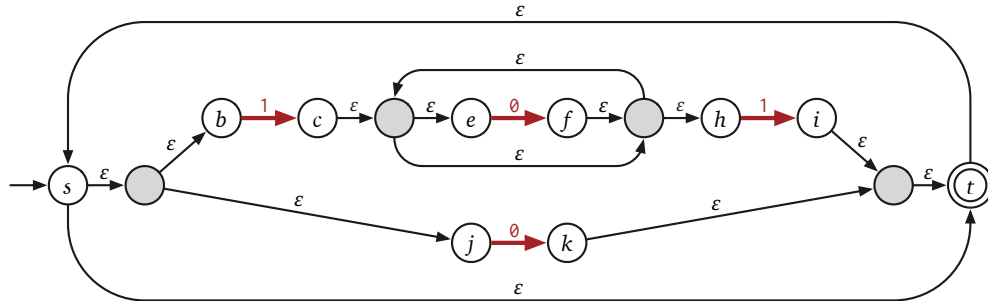
Interestingly, applying the incremental subset algorithm to Thompson's NFA tends to yield a DFA with relatively few states, in part because the states in Thompson's NFA tend to have large ϵ -reach, and in part because relatively few of those states are the sources of non- ϵ -transitions. Starting with the Thompson's NFA for $(0 + 10^*1)^*$, for example, the incremental subset construction yields a DFA with just five states.

⁵Unfortunately, the fastest NFA minimization algorithms run in exponential time, which is almost certainly inevitable, because the problem is PSPACE-hard. Also, unlike minimal DFAs, there can be multiple equivalent NFAs with the same minimum number of states.

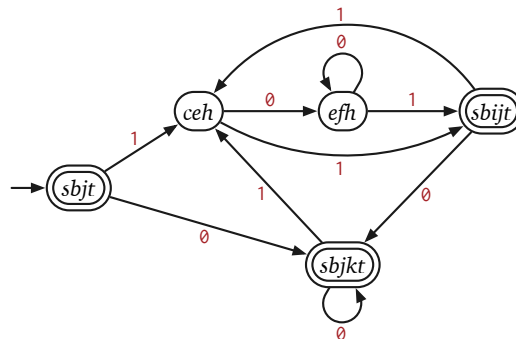
q'	$q' \in A'?$	$\delta'(q', 0)$	ϵ -reach	$\delta'(q', 1)$	ϵ -reach
ϵ -reach(s) = <i>sabjt</i>	✓	<i>k</i>	<i>sabjklt</i>	<i>c</i>	<i>cdegh</i>
<i>sabjklt</i>	✓	<i>k</i>	<i>sabjklt</i>	<i>c</i>	<i>cdegh</i>
<i>cdegh</i>		<i>f</i>	<i>defgh</i>	<i>i</i>	<i>sabijlt</i>
<i>defgh</i>		<i>f</i>	<i>defgh</i>	<i>i</i>	<i>sabijlt</i>
<i>sabijlt</i>	✓	<i>k</i>	<i>sabjklt</i>	<i>c</i>	<i>cdegh</i>



Keeping track of the large ϵ -reaches can be rather onerous if you're executing this algorithm by hand. We can save a bit of work by observing that in the incremental subset construction, we can ignore any states that are only connected to other states by ϵ -transitions, except for the start state s and the accepting state t .



q'	$q' \in A'?$	$\delta'(q', 0)$	ϵ -reach	$\delta'(q', 1)$	ϵ -reach
ϵ -reach(s) = <i>sbjt</i>	✓	<i>k</i>	<i>sbjkt</i>	<i>c</i>	<i>ceh</i>
<i>sbjkt</i>	✓	<i>k</i>	<i>sbjkt</i>	<i>c</i>	<i>ceh</i>
<i>ceh</i>		<i>f</i>	<i>efh</i>	<i>i</i>	<i>sbijt</i>
<i>efh</i>		<i>f</i>	<i>efh</i>	<i>i</i>	<i>sbijt</i>
<i>sbijt</i>	✓	<i>k</i>	<i>sbjkt</i>	<i>c</i>	<i>ceh</i>



This DFA can be further simplified to just two states, by observing that all three accepting states are equivalent, and that both non-accepting states are equivalent. But still, five states is pretty good, especially compared with the $2^{14} = 16384$ states (or 2^{10} interesting states) that the naïve subset construction would yield!

4.9 Another Example

On the next page is a series of figures showing another example of all the algorithms we've seen so far, starting with the regular expression $(0+1)^*(00+11)(0+1)^*$, which describes the language accepted by our very first example NFA. Thompson's algorithm constructs a 26-state monster.

The execution of the incremental subset construction and the resulting equivalent nine-state DFA are shown below the NFA. Yes, even if we ignore all the ε -only states, the ε -reaches get a bit ridiculous; unfortunately, this is typical for NFAs produced by Thompson's algorithm, which have lots of ε -transitions. (On the other hand, it's rare that anyone would execute this algorithm by hand. After all, that's why the gods gave us computers!) In the table, I've clustered the NFA-states *saceh* together, because those are part of *every* DFA-state, and I've clustered the states *kmt* together, because every DFA-state containing one of those NFA-states contains all three.

As usual, the resulting DFA still has far more states than necessary, but *significantly* less than the $2^{24} = 16\,777\,216$ states that the naïve subset construction would give us. (In the drawing of the DFAs, I've omitted the ε -reach NFA-states from the DFA-state names.)

Finally, the DFA-minimization algorithm from the previous lecture note correctly discovers that all six accepting states of the incremental-subset DFA are equivalent, and thus reduces the DFA to just four states, as shown on the next page.

*4.10 NFA to Regular Expression: State Elimination

The only component of Kleene's theorem left to prove is that every language accepted by an NFA is regular. I'll describe a proof that is morally equivalent to Kleene's 1951 argument, but uses more modern standard notation. This formulation of the proof was first published in 1960 by Robert McNaughton and Hisao Yamada⁶. A similar formulation (with much uglier notation) was published by C. Y. Lee earlier in the same year.⁷ The same algorithm is sometimes attributed to Janusz Brzozowski and Edward Mccluskey, who explored transformations of state machines in more generality a few years after McNaughton and Yamada.⁸

Recall that a standard NFA can be represented by a state-transition graph, whose vertices are the states and whose edges represent possible transitions. Each edge is labeled with a single symbol in Σ . A string $w \in \Sigma^*$ is accepted if and only if there is a sequence of transitions

$$s \xrightarrow{a_1} q_1 \xrightarrow{a_2} q_2 \xrightarrow{a_3} \cdots \xrightarrow{a_\ell} q_\ell$$

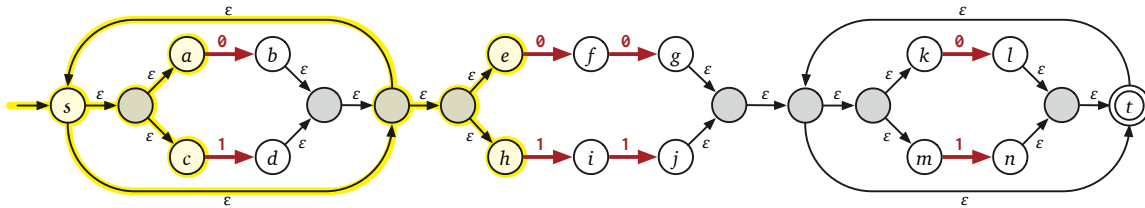
where the final state q_ℓ is accepting and $a_1 a_2 \cdots a_\ell = w$.

We've already seen that NFAs can be generalized to include ε -transitions; we can push this generalization further. A **string NFA** allows each transition $p \rightarrow q$ to be labeled with an *arbitrary*

⁶Robert McNaughton and Hisao Yamada. Regular expressions and state graphs for automata. *IRE Trans. Electron. Comput.* 9(1): 39-47, 1960.

⁷C. Y. Lee. Automata and finite automata. *Bell System Technical Journal* 39(2), 1267-1295, 1960.

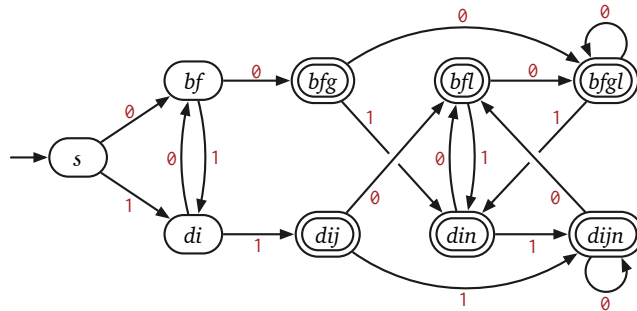
⁸Janusz A. Brzozowski and Edward J. McCluskey. Signal flow graph techniques for sequential circuit diagrams. *IEEE Trans. Electronic Computers*, EC-13(2):67-76, 1963.



Thompson's NFA for the regular expression $(0+1)^*(00+11)(0+1)^*$,
with the ϵ -reach $\{s, a, c, e, h\}$ of the start state s highlighted.

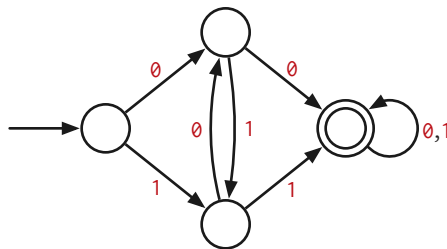
States that only connected to other states by ϵ -transitions are shaded and unnamed.

q'	$q' \in A'?$	$\delta'(q', 0)$	ϵ -reach	$\delta'(q', 1)$	ϵ -reach
ϵ -reach(s) = <i>saceh</i>		<i>bf</i>	<i>saceh bf</i>	<i>di</i>	<i>saceh di</i>
<i>saceh bf</i>		<i>bfg</i>	<i>saceh bfg kmt</i>	<i>di</i>	<i>saceh di</i>
<i>saceh di</i>		<i>bf</i>	<i>saceh bf</i>	<i>dij</i>	<i>saceh dij kmt</i>
<i>saceh bfg kmt</i>	✓	<i>bfgl</i>	<i>saceh bfgl kmt</i>	<i>din</i>	<i>saceh din kmt</i>
<i>saceh dij kmt</i>	✓	<i>bfl</i>	<i>saceh bfl kmt</i>	<i>dijn</i>	<i>saceh dijn kmt</i>
<i>saceh bfgl kmt</i>	✓	<i>bfgl</i>	<i>saceh bfgl kmt</i>	<i>din</i>	<i>saceh din kmt</i>
<i>saceh din kmt</i>	✓	<i>bfl</i>	<i>saceh bfl kmt</i>	<i>dijn</i>	<i>saceh dijn kmt</i>
<i>saceh bfl kmt</i>	✓	<i>bfgl</i>	<i>saceh bfgl kmt</i>	<i>din</i>	<i>saceh din kmt</i>
<i>saceh dijn kmt</i>	✓	<i>bfl</i>	<i>saceh bfl kmt</i>	<i>dijn</i>	<i>saceh dijn kmt</i>



The DFA computed by the incremental subset algorithm from Thompson's NFA for $(0+1)^*(00+11)(0+1)^*$.

The ϵ -reach NFA states *saceh* and *kmt* are omitted from the DFA-state names in the drawing.



The minimal DFA that accepts the language $(0+1)^*(00+11)(0+1)^*$.

string $x(p \rightarrow q) \in \Sigma^*$. We are allowed to transition from state p to state q if the label $x(p \rightarrow q)$ is a *prefix* of the remaining input. Thus, a string $w \in \Sigma^*$ is accepted if and only if there is a sequence of transitions

$$s \xrightarrow{x_1} q_1 \xrightarrow{x_2} q_2 \xrightarrow{x_3} \cdots \xrightarrow{x_\ell} q_\ell$$

where the final state q_ℓ is accepting, and $x_1 \cdot x_2 \cdot \cdots \cdot x_\ell = w$. Thus, an NFA with ε -transitions is just a string NFA where every label has length 0 or 1. Any string NFA can be converted into an equivalent standard NFA, by subdividing each edge $p \rightarrow q$ into a path of length $\max\{1, |x(p \rightarrow q)|\}$.

Finally, McNaughton and Yamada define an **expression automaton** as a finite-state machine where each transition $p \rightarrow q$ is labeled with an arbitrary *regular expression* $R(p \rightarrow q)$. We can transition from state p to state q if *any* prefix of the remaining input matches the regular expression $R(p \rightarrow q)$. Thus, a string $w \in \Sigma^*$ is accepted by an expression automaton if and only if there is a sequence of transitions

$$s \xrightarrow{R_1} q_1 \xrightarrow{R_2} q_2 \xrightarrow{R_3} \cdots \xrightarrow{R_\ell} q_\ell$$

where the final state q_ℓ is accepting, and a sequence of strings $x_1, x_2, \dots, x_\ell$ (any of which could be empty) such that $w = x_1 \cdot x_2 \cdot \cdots \cdot x_\ell$ and each substring x_i matches the corresponding regular expression R_i .

More formally, an expression automaton consists of the following five components:

- A finite set Σ called the **input alphabet**
- Another finite set Q whose elements are called **states**
- A unique **start state** $s \in Q$
- A unique **target state** $t \in Q \setminus \{s\}$
- A **transition function** $R: Q \times Q \rightarrow \text{Reg}(\Sigma)$, where $\text{Reg}(\Sigma)$ is the set of regular expressions over Σ .

The requirement that the target state is unique and distinct from s is not essential to the model, but imposed only to simplify the equivalence proof. Similarly, we assume for convenience that $R(q \rightarrow s) = R(t \rightarrow q) = \emptyset$ for every state q . Both of these assumptions can be enforced if necessary by adding artificial start and target states connected to other states by ε -transitions.

Expression automata are even more nondeterministic than NFAs. A single string could match several (even infinitely many) transition sequences from s to t , and it could match each of those sequences in several (even infinitely many) different ways. A string w is accepted if *any* decomposition of w into a sequence of substrings matches *any* sequence of transitions from s to t . Conversely, a string might match *no* state sequences, in which case the string is rejected.

Two extreme special cases of expression automata are already familiar. First, every regular language is clearly the language of an expression automaton with exactly two states. Second, with only minor modifications, any DFA or NFA can be converted into an expression automaton with trivial transition expressions. Thompson's algorithm can be used to transform any expression automaton into a standard NFA (with ε -transitions), by recursively expanding any nontrivial transition expression. To complete the proof of Kleene's theorem, we show how to convert an arbitrary expression automaton into a regular expression, by repeatedly deleting vertices.

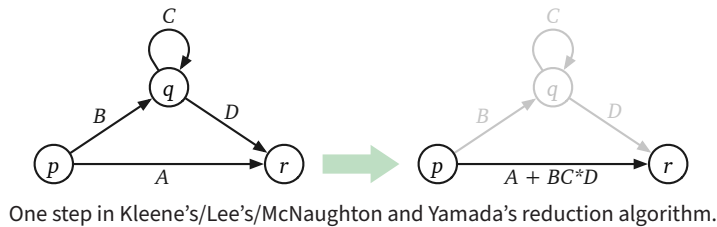
Lemma 4.3. *Every expression automaton accepts a regular language.*

Proof: Let $E = (Q, \Sigma, R, s, t)$ be an arbitrary expression automaton. Assume that any expression automaton with fewer states than E accepts a regular language. There are two cases to consider, depending on the number of states in Q :

- If $Q = \{s, t\}$, then E trivially accepts the regular language $R(s \rightarrow t)$.
- On the other hand, suppose Q has more than two states; fix an arbitrary state $q \in Q \setminus \{s, t\}$. We modify the automaton, without changing its language, so that state q is redundant and can be removed. Define a new transition function $R' : Q \times Q \rightarrow \text{Reg}(\Sigma)$ by setting

$$R'(p \rightarrow r) := R(p \rightarrow r) + R(p \rightarrow q)R(q \rightarrow q)^*R(q \rightarrow r).$$

With this modified transition function in place, any string w that matches the sequence $p \rightarrow q \rightarrow q \rightarrow \dots \rightarrow q \rightarrow r$, with any number of q 's, also matches the single transition $p \rightarrow r$. Thus, by induction, if w matches a sequence of states, it also matches the subsequence obtained by removing all q 's.



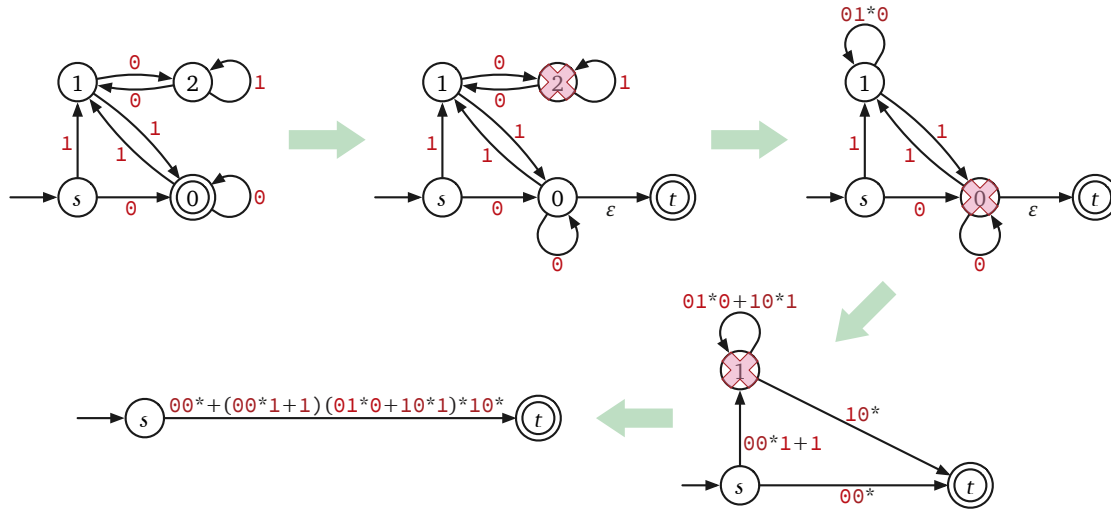
Let E' be the expression automaton with states $Q' = Q \setminus \{q\}$ that uses this modified transition function R' . This new automaton accepts exactly the same strings as the original automaton E . Because E' has fewer states than E , the inductive hypothesis implies E' accepts a regular language.

In both cases, we conclude that E accepts a regular language. $\square$

This proof can be mechanically translated into an algorithm to convert any DFA or NFA into an equivalent regular expression, via a sequence of expression automata with fewer and fewer states, but increasingly complex transition expressions.

The figure below shows McNaughton and Yamada's reduction algorithm in action, starting with a DFA that accepts the binary representations of non-negative integers divisible by 3, possibly with extra leading 0s. (State i means the binary number we've read so far is congruent to $i \bmod 3$.) We add an artificial target state t with an ε -transition from the unique accepting state 0. (We don't need an artificial start state, because s already has no incoming transitions.) We then remove the numbered states one by one, updating the transition expressions between any remaining states at each iteration. For the sake of clarity, the figure omits all edges $p \rightarrow q$ where $R(p \rightarrow q) = \emptyset$. The final regular expression $0^*0 + (00^*1 + 1)(10^*1 + 01^*0)^*10^*$ can be slightly simplified to $0^*0 + 0^*1(10^*1 + 01^*0)^*10^*$, which is precisely the regular expression we gave for this language back in Lecture Note 2!

Given an NFA with n states (including s and t), McNaughton and Yamada's algorithm iteratively removes $n - 2$ states, updating $O(n^2)$ transition expressions in each iteration. If the concatenation and Kleene star operations could be performed in constant time, the resulting algorithm would run in $O(n^3)$ time. However, in the worst case, the transition expressions grow



Converting a DFA into an equivalent regular expression using McNaughton and Yamada's algorithm.

in length by roughly a factor of 4 in each iteration; the final expression has length $\Theta(4^n)$. If we represent the transition expressions either as explicit strings or as regular expression trees, the worst-case running time is also $\Theta(4^n)$.

4.11 Regular Language Transformations

We have already seen that many functions of regular languages are themselves regular: unions, concatenations, and Kleene closure by definition. and intersections and differences by product construction on DFAs. However, the set of regular languages is closed under a much richer class of functions.

Suppose we wanted to prove that regular languages are closed under some function f ; that is, for every regular language L , we want to prove that the language $f(L)$ is also regular. There are two general techniques to prove such a statement:

- Describe an algorithm that transforms an arbitrary regular expression R into a new regular expression R' such that $L(R') = f(L(R))$.
- Describe an algorithm that transforms an arbitrary DFA M into a new NFA M' such that $L(M') = f(L(M))$.

The equivalence between regular expressions and finite automata implies that *in principle* we can always use either technique, but in practice, the second one is far more powerful and usually simpler. The asymmetry in the second technique is important. We start with a DFA for L to impose as much structure as possible in the input; we aim for an NFA with ε -transitions to give ourselves as much freedom as possible in the output.⁹

4.11.1 Reversal

For our first example, I'll describe proofs using both techniques.

⁹We could give ourselves even more freedom by constructing an *expression automaton*, or even better, a *generalized expression automaton*, which would allow intersections and negations in the transition expressions. On the other hand, creativity thrives on constraint.

Lemma 4.4. *For any regular language L , the language $L^R = \{w^R \mid w \in L\}$ is also regular.*

Proof (regular expression to regular expression): Let R be an arbitrary regular expression such that $L = L(R)$. Assume for any proper subexpression S of R that $L(S)^R$ is regular. There are five cases to consider, mirroring the recursive definition of regular expressions:

- If $R = \emptyset$, then $L^R = L = \emptyset$, so $L(R) = L^R$.
- Suppose R consists of a single word w . Let $R' = w^R$. Then $L(R') = \{w^R\} = L^R$.
- Suppose $R = A + B$. The inductive hypothesis implies that there are regular expressions A' and B' such that $L(A') = L(A)^R$ and $L(B') = L(B)^R$. Let $R' = A' + B'$. Then $L(R') = L(A') \cup L(B') = L(A)^R \cup L(B)^R = (L(A) \cup L(B))^R = L^R$.
- Suppose $R = A \cdot B$. The inductive hypothesis implies that there are regular expressions A' and B' such that $L(A') = L(A)^R$ and $L(B') = L(B)^R$. Let $R' = B' \cdot A'$. Then $L(R') = L(B') \cdot L(A') = L(B)^R \cdot L(A)^R = (L(A) \cdot L(B))^R = L^R$.
- Finally, suppose $R = A^*$. The inductive hypothesis implies that there is a regular expression A' such that $L(A') = L(A)^R$. Let $R' = (A')^*$. Then $L(R') = L(A')^* = (L(A)^R)^* = (L(A)^*)^R = L^R$.

In all cases, we have constructed a regular expression R' such that $L(R') = L^R$. We conclude that L^R is regular. $\square$

For example, if L is described by the regular expression $(10 + 1)^*(01 + 100)$, then L^R is described by the regular expression $(10 + 001)(0 + 01)^*$.

Careful readers may be unsatisfied with the previous proof, because it assumes several “obvious” properties of string and language reversal. Specifically, for all strings x and y and all languages L and L' , we assumed the following identities:

- $(x \cdot y)^R = y^R \cdot x^R$
- $(L \cdot L')^R = (L')^R \cdot L^R$
- $(L \cup L')^R = L^R \cup (L')^R$
- $(L^*)^R = (L^R)^*$

All of these claims are all easy to prove by inductive definition-chasing.

However, in general, it is much easier to argue regularity of the transformed language by constructing an NFA.

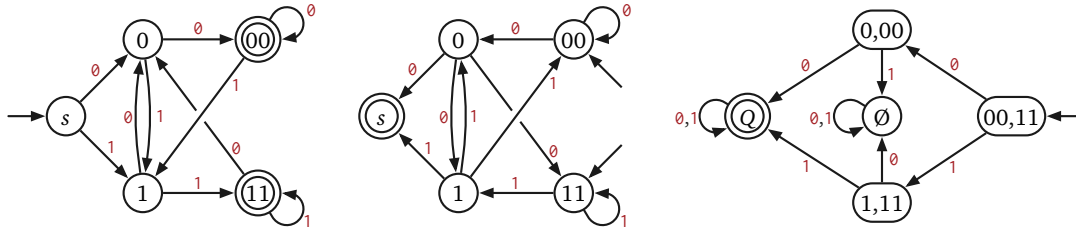
Proof (DFA to NFA): Let $M = (\Sigma, Q, s, A, \delta)$ be an arbitrary DFA that accepts L . We construct an NFA $M^R = (\Sigma, Q^R, S^R, A^R, \delta^R)$ that accepts L^R , intuitively by reversing every transition in M , and swapping the roles of the start state and the accepting states. Because M may not have a unique accepting state, the reversed NFA M' has multiple start states.

$$\begin{aligned} Q^R &= Q \\ S^R &= A \\ A^R &= \{s\} \\ \delta^R(q, a) &= \{p \mid q = \delta(p, a)\} \end{aligned} \quad \text{for all } q \in Q \text{ and } a \in \Sigma$$

Routine inductive definition-chasing now implies that the reversal of any sequence $s \rightarrow q_1 \rightarrow \dots \rightarrow q_\ell$ of transitions in M is a valid sequence $q_\ell \rightarrow \dots \rightarrow q_1 \rightarrow s$ of transitions in M^R and vice versa. Because the transitions retain their labels (but reverse directions), it follows that M accepts any string w if and only if M^R accepts w^R .

We conclude that the NFA M^R accepts L^R , and therefore L^R must be regular. $\square$

The following figure shows an example of this transformation in action. On the left is a DFA that accepts binary strings whose last two symbols are equal. Applying our reversal transformation yields the NFA in the middle. Finally, applying the incremental subset construction to this reversed NFA yields the DFA on the right. The final DFA clearly accepts binary strings whose *first* two symbols are equal, and therefore so does the equivalent middle NFA.



Reversing a DFA, and then transforming the resulting NFA into an equivalent DFA.

4.11.2 First Half of Palindromes

Lemma 4.5. For any regular language L , the language $\text{palin}(L) := \{w \mid ww^R \in L\}$ is also regular.

The language $\text{palin}(L)$ contains the first half of every even-length palindrome in L . For example, if $L = \{\varepsilon, \text{AHA}, \text{POOP}, \text{TACOCAT}, \text{PALINDROME}\}$, then $\text{palin}(L) = \{\varepsilon, \text{PO}\}$.

As a more interesting example, suppose L is the language $(0 + 1)^*(00 + 11)$ of strings whose last two symbols are equal. Every palindrome in L either starts and ends with 00 , or starts and ends with 11 . Thus, $\text{palin}(L)$ is actually the language $0 + 1 + (00 + 11)(0 + 1)^*$ of strings either with length one or whose first two symbols are equal.

Proof: Let $M = (\Sigma, Q, s, A, \delta)$ be an arbitrary DFA that accepts L .

Intuitively, we construct an NFA M' that reads its input string w and simulates the original DFA M reading the input string ww^R . The NFA M' is a product construction of the original DFA M and its reversal NFA M^R constructed in our second proof of Lemma 4.4. Just like any other product construction, we must explicitly specify the accepting states to complete the description of M' .

Let $q = \delta^*(s, w)$ be the state that M reaches on reading input string w . Then M accepts the palindrome ww^R if and only if

$$\delta^*(s, ww^R) = \delta^*(\delta^*(s, w), w^R) = \delta^*(q, w^R) \in A.$$

But $\delta^*(q, w^R) \in A$ if and only if the reversed machine M^R reaches state q on reading the same input string w . Thus, our new NFA M' should accept w if and only if both M and M^R end at the same state. In other words, we need $A' = \{(q, q) \mid q \in Q\}$.

Summing up, we construct an NFA $M' = (\Sigma, Q', S', A', \delta')$ with multiple start states as follows:

$$\begin{aligned} Q' &= Q \times Q \\ S' &= \{(s, q) \mid q \in A\} \\ A' &= \{(q, q) \mid q \in Q\} \\ \delta'((q, r), a) &= \{(\delta(q, a), p) \mid r = \delta(p, a)\} \quad \text{for all } q, r \in Q \text{ and } a \in \Sigma \quad \square \end{aligned}$$

The NFAs constructed by this algorithm can be quite large. For example, if M is the five-state DFA for the language $L = (0 + 1)^*(00 + 11)$ that we used to illustrate the reversal transformation, our algorithm constructs an NFA M' for $\text{palin}(L)$ with 30 states, which is far too complex to draw. It turns out that 11 of those 30 states either cannot be reached from any start state (s, q) or cannot reach an accepting state (q, q) ; more significantly, the language $\text{palin}(L) = 0 + 1 + (00 + 11)(0 + 1)^*$ can actually be decided by a five-state DFA! But that extra complexity from our algorithm is completely okay—we're not even trying to construct an *efficient* machine; we just need to prove that one NFA exists.

4.11.3 A Harder Example

In our last example, we make more fundamental use of nondeterminism to *guess* the future behavior of the DFA we are simulating.

Lemma 4.6. *For any regular language L , the language $\text{half}(L) := \{w \mid ww \in L\}$ is also regular.*

Proof: Let $M = (\Sigma, Q, s, A, \delta)$ be an arbitrary DFA that accepts L .

Intuitively, we construct an NFA M' that reads its input string w and simulates the original DFA M reading the input string ww . Our overall strategy has three parts:

- First M' non-deterministically *guesses* the state $h = \delta^*(s, w)$ that M will reach after reading input w . (We can't just run M on input w to compute the correct state h , because that would consume the input string!)
- Then M' runs two copies of M in parallel (using a product construction): a “left” copy starting at s and a “right” copy starting at the (guessed) halfway state h .
- Finally, when M' is done reading w , it accepts if and only if the first copy of M actually stopped in state h (so our initial guess was correct) and the second copy of M stopped in an accepting state. That is, M' accepts if and only if $\delta^*(s, w) = h$ and $\delta^*(h, w) \in A$.

To implement this strategy, M' needs to maintain *three* states of M : the state of the left copy of M , the guess h for the halfway state, and the state of the right copy of M . The first and third states evolve according to the transition function δ , but the second state never changes. Finally, to implement the non-deterministic guessing, M' has multiple start states.

Summing up, our new NFA $M' = (\Sigma, Q', S', A', \delta')$ is formally defined as follows.

$$\begin{aligned} Q' &= Q \times Q \times Q \\ S' &= \{(s, h, h) \mid h \in S\} \\ A' &= \{(h, h, q) \mid h \in Q \text{ and } q \in A\} \\ \delta'((p, h, q), a) &= \{(\delta(p, a), h, \delta(q, a))\} \quad \text{for all } p, h, q \in Q \text{ and } a \in \Sigma \quad \square \end{aligned}$$

Exercises

1. For each of the following regular expressions, describe or draw two finite-state machines:

- An NFA that accepts the same language, constructed using Thompson's algorithm.
- An equivalent DFA, built from the previous NFA using the incremental subset construction. For each state in your DFA, identify the corresponding subset of states in your NFA. Your DFA should have no unreachable states.

- (a) $(01 + 10)^*(0 + 1 + \varepsilon)$
 (b) $(\varepsilon + 1)(01)^*(\varepsilon + 0)$
 (c) $1^* + (10)^* + (100)^*$
 (d) $(\varepsilon + 0 + 00)(1 + 10 + 100)^*$
 (e) $((0 + 1)(0 + 1))^*$
 (f) $\varepsilon + 0(0 + 1)^* + 1(1 + 0)^*$

2. The accepting language of an NFA $M = (\Sigma, Q, s, A, \delta)$ is defined as follows:

$$L(M) := \{w \in \Sigma^* \mid \delta^*(s, w) \cap A \neq \emptyset\}.$$

Kleene's theorem (described here as McNaughton and Yamada's state-reduction algorithm) implies that $L(M)$ is regular. Prove that the following languages associated with M are also regular:

- (a) $L^\forall(M) := \{w \in \Sigma^* \mid A \subseteq \delta^*(s, w)\}$. That is, a string w is in the language $L^\forall(M)$ if and only if $\delta^*(s, w)$ contains *every* accepting state.
 (b) $L^\subseteq(M) := \{w \in \Sigma^* \mid \delta^*(s, w) \subseteq A\}$. That is, a string w is in the language $L^\subseteq(M)$ if and only if $\delta^*(s, w)$ contains *only* accepting states.
 (c) $L^\equiv(M) := \{w \in \Sigma^* \mid \delta^*(s, w) = A\}$. That is, a string w is in the language $L^\equiv(M)$ if and only if $\delta^*(s, w)$ is exactly the set of accepting states.

[Hint: These are easier than they look!]

3. Let $M = (Q, S, A, \delta)$ be an NFA, possibly with ε -transitions and multiple start states. Given any input string w , the NFA nondeterministically chooses a sequence of transitions $s \xrightarrow{a_1} q_1 \xrightarrow{a_2} q_2 \xrightarrow{a_3} \dots \xrightarrow{a_\ell} q_\ell$ where each a_i is either ε or a symbol in Σ and $a_1 a_2 \dots a_\ell = w$. We call any such sequence a **computation path** for w in M . We call a computation path is *accepting* if it ends in an accepting state and *rejecting* otherwise.

The standard definition of NFA acceptance is that M accepts w if and only if there is *at least one* accepting computation path for w in M . This questions asks you to consider alternative acceptance criteria.

- (a) A *universal* finite-state automaton, abbreviated **$\forall$ FA**, accepts w if and only if *every* computation path for w is accepting, or equivalently, if there are no rejecting computation paths for w .

- i. Prove that the set of all strings accepted by a $\forall$ FA is regular.
 - ii. Prove that every regular language is the set of strings accepted by a $\forall$ FA.
 - ★(b) A *parity* finite-state automaton, abbreviated $\oplus$ FA, is an NFA M **without ε -transitions** that accepts w if and only if there are *an odd number of* accepting computation paths for w in M .
 - i. Prove that the set of all strings accepted by a $\oplus$ FA is regular.
 - ii. Prove that every regular language is the set of strings accepted by a $\oplus$ FA.
 - ★(c) An *infinitary* finite-state automaton, abbreviated ∞ FA, accepts w if and only if there are *infinitely many* accepting computation paths for w .
 - i. Prove that the set of all strings accepted by an ∞ FA is regular.
 - ii. Prove that every regular language is the set of strings accepted by an ∞ FA.

[Hint: If there are infinitely many computation paths for w , there must be a cycle of ε -transitions in M .]
 - ★(d) A *blind* infinitary finite-state automaton, abbreviated $B\infty$ FA, accepts w if and only if there are infinitely many computation paths for w . Notice that accepting states play absolutely no role in this definition!
 - i. Prove that the set of all strings accepted by a $B\infty$ FA is regular.
 - ii. Prove that every regular language is the set of strings accepted by a $B\infty$ FA.
4. A certain professor who really should know better once woke up in the middle of the night with a startling revelation—Thompson's algorithm doesn't need all those ε -transitions! Filled with the certainty that only sleep deprivation can bring, he ran to his laptop and quickly changed two cases in his description of Thompson's algorithm.

- When $R = S \cdot T$, instead of connecting the accept state of $\text{NFA}(S)$ to the start state of $\text{NFA}(T)$ with an ε -transition, we can **just** identify those two states to build the simpler NFA $\text{NFA}(S \cdot T)$!
- When $R = S^*$, instead of introducing two new states and four ε -transitions, we can **just** add two ε -transitions between the start and accept states of $\text{NFA}(S)$ to build the



Satisfied with his simplification, he thanked the penguin who gave him the idea, and then flew his hat back into the ocean marshmallows, where a giant man with the head of a dog gave him the power of bread.¹⁰ The next morning, while he was proudly teaching his new simplified proof for the first time, he realized his embarrassing mistake.

Prove that *neither* of the professor's optimizations is actually correct.

- (a) Find a regular expression R , such that the NFA constructed from R by Thompson's algorithm **with only the first modification** accepts strings that are not in $L(R)$.

¹⁰Eid ma clack shaw supoven del ba. Mertepty ven seinur cofally ragdah.

- (b) Find a regular expression R , such that the NFA constructed from R by Thompson's algorithm **with only the second modification** accepts strings that are not in $L(R)$.

The moral of this story: Whenever you use the word “**just**”, you're probably missing something important.

5. A finite-state *transducer* is a finite-state machine that produces output as it reads the input. The simplest class of finite-state transducers are called **Moore machines**. For purposes of this problem, a Moore machine formally consists of six components:

- A finite set Σ called the input alphabet
- A finite set Γ called the output alphabet
- A finite set Q whose elements are called states
- A start state $s \in Q$
- A transition function $\delta: Q \times \Sigma \rightarrow Q$
- An output function $\omega: Q \rightarrow \Gamma$

More intuitively, a Moore machine is a graph with a special start vertex, where every node (state) has one outgoing edge labeled with each symbol from the input alphabet, and each node (state) is additionally labeled with a symbol from the output alphabet.

Each Moore machine reads an input string $w \in \Sigma^*$ one symbol at a time and produces a single output symbol for each symbol it reads. Specifically, if the machine is in state q and reads symbol a , it changes its state to $q' = \delta(q, a)$ and then outputs the symbol $\omega(q')$. Formally, we recursively define a *transducer* function $\omega^*: Q \times \Sigma^* \rightarrow \Gamma^*$ as follows:

$$\omega^*(q, w) = \begin{cases} \varepsilon & \text{if } w = \varepsilon \\ \omega(\delta(q, a)) \cdot \omega^*(\delta(q, a), x) & \text{if } w = ax \end{cases}$$

Given the input string $w \in \Sigma^*$, the machine outputs the string $\omega^*(s, w) \in \Gamma^*$. To simplify notation, we define $M(w) = \omega^*(s, w)$.

- (a) Let M be an arbitrary Moore machine. The **output language** $L^\circ(M)$ of a Moore machine M is the set of all strings that the machine can output:

$$L^\circ(M) := \{M(w) \mid w \in \Sigma^*\}$$

Prove that $L^\circ(M)$ is a regular language.

- (b) Let M be an arbitrary Moore machine whose input alphabet Σ and output alphabet Γ are identical. Prove that the language

$$L^=(M) = \{w \in \Sigma^* \mid M(w) = w\}$$

is regular. Strings in $L^=(M)$ are also called *fixed points* of the function $M: \Sigma^* \rightarrow \Sigma^*$.

- * (c) As in part (b), let M be an arbitrary Moore machine whose input and output alphabets are identical. Prove that the language $\{w \in \Sigma^* \mid M(M(w)) = w\}$ is regular.

[Hint: Parts (a) and (b) are easier than they look!]

6. Another slightly more general class of finite-state transducers are called **Mealy machines**. Mealy machines are similar to the Moore machines considered in the previous problem, but they produce outputs at each *transition* instead of at each *state*.

For purposes of this problem, a Mealy machine formally consists of six components:

- A finite set Σ called the input alphabet
- A finite set Γ called the output alphabet
- A finite set Q whose elements are called states
- A start state $s \in Q$
- A transition function $\delta: Q \times \Sigma \rightarrow Q$
- An output function $\omega: Q \times \Sigma \rightarrow \Gamma$

More intuitively, a Mealy machine is a graph with a special start vertex, where every node (state) has one outgoing edge labeled with each symbol from the input alphabet, and each edge (transition) is additionally labeled with a symbol from the output alphabet.

As usual, each Mealy machine reads its input string $w \in \Sigma^*$ one symbol at a time and produces a single output symbol for each symbol it reads. Specifically, if the machine is in state q and reads symbol a , it outputs the symbol $\omega(q, a)$ and then changes its state to $\delta(q, a)$. Formally, we recursively define a *transducer function* $\omega^*: Q \times \Sigma^* \rightarrow \Gamma^*$ as follows:

$$\omega^*(q, w) = \begin{cases} \varepsilon & \text{if } w = \varepsilon \\ \omega(q, a) \cdot \omega^*(\delta(q, a), x) & \text{if } w = ax \end{cases}$$

Given any input string $w \in \Sigma^*$, the machine outputs the string $\omega^*(s, w) \in \Gamma^*$. To simplify notation, we define $M(w) = \omega^*(s, w)$.

- (a) Let M be an arbitrary **Moore** machine. Describe a **Mealy** machine that defines exactly the same transducer function ω^* .
- (b) Let M be an arbitrary Mealy machine. The **output language** $L^\circ(M)$ of a Mealy machine M is the set of all strings that the machine can output:

$$L^\circ(M) := \{M(w) \mid w \in \Sigma^*\}$$

Prove that $L^\circ(M)$ is a regular language.

- (c) Let M be an arbitrary Mealy machine whose input alphabet Σ and output alphabet Γ are identical. Prove that the language

$$L^-(M) = \{w \in \Sigma^* \mid w = \omega^*(s, w)\}$$

is regular. $L^-(M)$ consists of all strings w such that M outputs w when given input w ; these are also called *fixed points* for the transducer function ω^* .

- * (d) As in part (c), let M be an arbitrary Mealy machine whose input and output alphabets are identical. Prove that the language $\{w \in \Sigma^* \mid M(M(w)) = w\}$ is regular.

[Hint: Parts (a), (b), and (c) are easier than they look!]

7. Let L be an arbitrary regular language, over an arbitrary alphabet Σ that contains the symbol $\#$. Prove that the following languages are regular.

- (a) $\text{censor}(L) := \{\#^{|w|} \mid w \in L\}$
- (b) $\text{dehash}(L) = \{\text{dehash}(w) \mid w \in L\}$, where $\text{dehash}(w)$ is the subsequence of w obtained by deleting every $\#$.
- (c) $\text{insert}\#(L) := \{x\#y \mid xy \in L\}$.
- (d) $\text{delete}\#(L) := \{xy \mid x\#y \in L\}$.
- (e) $\text{prefix}(L) := \{x \in \Sigma^* \mid xy \in L \text{ for some } y \in \Sigma^*\}$
- (f) $\text{suffix}(L) := \{y \in \Sigma^* \mid xy \in L \text{ for some } x \in \Sigma^*\}$
- (g) $\text{substring}(L) := \{y \in \Sigma^* \mid xyz \in L \text{ for some } x, z \in \Sigma^*\}$
- (h) $\text{superstring}(L) := \{xyz \mid y \in L \text{ and } x, z \in \Sigma^*\}$
- (i) $\text{cycle}(L) := \{xy \mid x, y \in \Sigma^* \text{ and } yx \in L\}$
- (j) $\text{prefmax}(L) := \{x \in L \mid xy \in L \iff y = \varepsilon\}$.
- (k) $\text{sufmin}(L) := \{xy \in L \mid y \in L \iff x = \varepsilon\}$.
- (l) $\text{minimal}(L) := \{w \in L \mid \text{no proper substring of } w \text{ is in } L\}$.
- (m) $\text{maximal}(L) := \{w \in L \mid \text{no proper superstring of } w \text{ is in } L\}$.
- (n) $\text{evens}(L) := \{\text{evens}(w) \mid w \in L\}$, where $\text{even}(w)$ is the subsequence of w containing every even-indexed symbol. For example, $\text{evens}(\text{EVENINDEX}) = \text{VNNE}$.
- (o) $\text{evens}^{-1}(L) := \{w \in \Sigma^* \mid \text{evens}(w) \in L\}$.
- (p) $\text{subseq}(L) := \{x \in \Sigma^* \mid x \text{ is a subsequence of some } y \in L\}$
- (q) $\text{superseq}(L) := \{x \in \Sigma^* \mid \text{some } y \in L \text{ is a subsequence of } x\}$
- (r) $\text{swap}(L) := \{\text{swap}(w) \mid w \in L\}$, where $\text{swap}(w)$ is defined recursively as follows:

$$\text{swap}(w) = \begin{cases} w & \text{if } |w| \leq 1 \\ ba \cdot \text{swap}(x) & \text{if } w = abx \text{ for some } a, b \in \Sigma \text{ and } x \in \Sigma^* \end{cases}$$

- (s) $\text{oneswap}(L) := \{xbay \mid xaby \in L \text{ where } a, b \in \Sigma \text{ and } x, y \in \Sigma^*\}$.

8. Let L be an arbitrary regular language over an arbitrary alphabet Σ . Prove that the following languages are regular.

- (a) $\text{left}(L) := \{x \in \Sigma^* \mid xy \in L \text{ for some } y \in \Sigma^* \text{ where } |x| = |y|\}$
- (b) $\text{right}(L) := \{y \in \Sigma^* \mid xy \in L \text{ for some } x \in \Sigma^* \text{ where } |x| = |y|\}$

- (c) $middle(L) := \{y \in \Sigma^* \mid xyz \in L \text{ for some } x, z \in \Sigma^* \text{ where } |x| = |y| = |z|\}$
- (d) $halfseq(L) := \{w \in \Sigma^* \mid w \text{ is a subsequence of some string } x \in L \text{ where } |x| = 2 \cdot |w|\}$
- (e) $third(L) := \{w \in \Sigma^* \mid www \in L\}$
- (f) $drome(L) := \{w \in \Sigma^* \mid w^R w \in L\}$

9. Let L and L' be arbitrary regular languages over the alphabet $\{0, 1\}$. Prove that the following languages are also regular:

- (a) $L \sqcap L' := \{x \sqcap y \mid x \in L \text{ and } y \in L' \text{ and } |x| = |y|\}$, where $x \sqcap y$ denotes bitwise-and. For example, $0011 \sqcap 0101 = 0001$.
- (b) $L \sqcup L' := \{x \sqcup y \mid x \in L \text{ and } y \in L' \text{ with } |x| = |y|\}$, where $x \sqcup y$ denotes bitwise-or. For example, $0011 \sqcup 0101 = 0111$.
- (c) $L \boxplus L' := \{x \boxplus y \mid x \in L \text{ and } y \in L' \text{ with } |x| = |y|\}$, where $x \boxplus y$ denotes bitwise-exclusive-or. For example, $0011 \boxplus 0101 = 0110$.
- (d) $faro(L, L') := \{faro(x, z) \mid x \in L \text{ and } z \in L' \text{ with } |x| = |z|\}$, where

$$faro(x, z) := \begin{cases} z & \text{if } x = \varepsilon \\ a \cdot faro(z, y) & \text{if } x = ay \end{cases}$$

For example, $faro(0011, 0101) = 00011011$.

- (e) $shuffles(L, L') := \bigcup_{w \in L, y \in L'} shuffles(w, y)$, where $shuffles(w, y)$ is the set of all strings obtained by shuffling w and y , or equivalently, all strings in which w and y are complementary subsequences. Formally:

$$shuffles(w, y) = \begin{cases} \{y\} & \text{if } w = \varepsilon \\ \{w\} & \text{if } y = \varepsilon \\ \{a\} \cdot shuffles(x, y) \cup \{b\} \cdot shuffles(w, z) & \text{if } w = ax \text{ and } y = bz \end{cases}$$

For example, $shuffles(01, 10) = \{0101, 0110, 1001, 1010\}$ and $shuffles(00, 11) = \{0011, 0101, 1001, 0110, 1010, 1100\}$.

10. (a) Let $inc: \{0, 1\}^* \rightarrow \{0, 1\}^*$ denote the *increment* function, which transforms the binary representation of an arbitrary integer n into the binary representation of $n + 1$, truncated to the same number of bits. For example:

$$inc(0010) = 0011 \quad inc(0111) = 1000 \quad inc(1111) = 0000 \quad inc(\varepsilon) = \varepsilon$$

Let $L \subseteq \{0, 1\}^*$ be an arbitrary regular language. Prove that $inc(L) = \{inc(w) \mid w \in L\}$ is also regular.

- (b) Let $dec: \{0, 1\}^* \rightarrow \{0, 1\}^*$ denote the *decrement* function, which transforms the binary representation of an arbitrary integer n into the binary representation of $n - 1$,

truncated to the same number of bits. More simply, dec is the inverse function of inc . For example:

$$dec(0011) = 0010 \quad dec(1000) = 0111 \quad dec(0000) = 1111 \quad dec(\varepsilon) = \varepsilon$$

Let $L \subseteq \{0, 1\}^*$ be an arbitrary regular language. Prove that $dec(L) = \{dec(w) \mid w \in L\}$ is also regular.

- (c) Let $dub: \{0, 1\}^* \rightarrow \{0, 1\}^*$ denote the *doubling* function, which transforms the binary representation of an arbitrary integer n into the binary representation of $2n$, truncated to the same number of bits. For example:

$$dub(0010) = 0100 \quad dub(0111) = 1110 \quad dub(1111) = 1110 \quad dub(\varepsilon) = \varepsilon$$

Let $L \subseteq \{0, 1\}^*$ be an arbitrary regular language. Prove that $dub(L) = \{dub(w) \mid w \in L\}$ is also regular.

- * (d) Let $trip: \{0, 1\}^* \rightarrow \{0, 1\}^*$ denote the *tripling* function, which transforms the binary representation of an arbitrary integer n into the binary representation of $3n$, truncated to the same number of bits. For example:

$$trip(0010) = 0110 \quad trip(0111) = 0101 \quad trip(1111) = 1101 \quad trip(\varepsilon) = \varepsilon$$

Let $L \subseteq \{0, 1\}^*$ be an arbitrary regular language. Prove that $trip(L) = \{trip(w) \mid w \in L\}$ is also regular. [Hint: It might be easier to consider the language $trip(L^R)^R$ first. Then again, it might not.]

- *11. Let $L \subseteq \Sigma^*$ be an arbitrary regular language. Prove that the following languages are regular.

- (a) $\text{sqrt}(L) := \{x \in \Sigma^* \mid xy \in L \text{ for some } y \in \Sigma^* \text{ such that } |y| = |x|^2\}$
- (b) $\text{log}(L) := \{x \in \Sigma^* \mid xy \in L \text{ for some } y \in \Sigma^* \text{ such that } |y| = 2^{|x|}\}$
- (c) $\text{flog}(L) := \{x \in \Sigma^* \mid xy \in L \text{ for some } y \in \Sigma^* \text{ such that } |y| = F_{|x|}\}$, where F_n is the n th Fibonacci number.

- *12. Let $L \subseteq \Sigma^*$ be an arbitrary regular language. Prove that the following languages are regular.

- (a) $\text{somerep}(L) := \{w \in \Sigma^* \mid w^n \in L \text{ for some } n \geq 0\}$
- (b) $\text{allreps}(L) := \{w \in \Sigma^* \mid w^n \in L \text{ for every } n \geq 0\}$
- (c) $\text{manyreps}(L) := \{w \in \Sigma^* \mid w^n \in L \text{ for infinitely many } n \geq 0\}$
- (d) $\text{fewreps}(L) := \{w \in \Sigma^* \mid w^n \in L \text{ for finitely many } n \geq 0\}$
- (e) $\text{powers}(L) := \{w \in \Sigma^* \mid w^{2^n} \in L \text{ for some } n \geq 0\}$
- ★ (f) $\text{whatthe}_N(L) := \{w \in \Sigma^* \mid w^n \in L \text{ for some } n \in N\}$, where N is an **arbitrary** fixed set of non-negative integers. [Hint: You only have to prove that an accepting NFA exists; you don't have to describe how to construct it.]

[Hint: For each of these languages, there is an accepting NFA with at most q^q states, where q is the number of states in some DFA that accepts L .]

13. For any string $w \in (\mathbf{0} + \mathbf{1})^*$, let $\text{sort}(w)$ be the string obtained by sorting the symbols in w . More succinctly, let $\text{sort}(w) = \mathbf{0}^{\#(\mathbf{0},w)}\mathbf{1}^{\#(\mathbf{1},w)}$.
- (a) Prove that there is a regular language $L \subseteq (\mathbf{0} + \mathbf{1})^*$ such that the language $\text{sort}(L) = \{\text{sort}(w) \mid w \in L\}$ is *not* regular.
- * (b) Prove that for any regular language $L \subseteq (\mathbf{0} + \mathbf{1})^*$, the language $\text{sort}^{-1}(L) = \{w \mid \text{sort}(w) \in L\}$ is regular.

- *14. For any string $w \in (\mathbf{0} + \mathbf{1})^*$, let $\langle w \rangle_2$ denote the integer represented by w in binary. For example:

$$\langle \varepsilon \rangle_2 = 0 \quad \langle \mathbf{0010} \rangle_2 = 2 \quad \langle \mathbf{0111} \rangle_2 = 7 \quad \langle \mathbf{1111} \rangle_2 = 15$$

Let L and L' be arbitrary regular languages over the alphabet $\{\mathbf{0}, \mathbf{1}\}$. Prove that the following language is also regular:

$$\{w \in (\mathbf{0} + \mathbf{1})^* \mid \langle w \rangle_2 = \langle x \rangle_2 + \langle y \rangle_2 \text{ for some strings } x \in L \text{ and } y \in L'\}$$

- *15. Let $L \subseteq \Sigma^*$ be an arbitrary regular language. Prove that the following languages are regular.

- (a) $\text{repsqrt}(L) = \{w \in \Sigma^* \mid w^{|w|} \in L\}$.
- (b) $\text{replog}(L) = \{w \in \Sigma^* \mid w^{2^{|w|}} \in L\}$.
- (c) $\text{repflog}(L) = \{w \in \Sigma^* \mid w^{F_{|w|}} \in L\}$, where F_n is the n th Fibonacci number.

[Hint: The NFAs for these languages use a **LOT** of states. Let $M = (\Sigma, Q, s, A, \delta)$ be a DFA that accepts L . Imagine that you somehow know $\delta^*(q, w)$ in advance, for every state $q \in Q$. Ha, ha, ha! Mine is an evil laugh!]

- *16. A **nondeterministic finite-state transducer** or *NFT* is a generalization of the Mealy machine—a finite-state machine that produces output. Formally, an NFT consists of the following components:

- A non-empty finite *input* alphabet Σ
- A non-empty finite *output* alphabet Λ
- A directed graph $G = (Q, E)$, possibly with loops and parallel edges, whose vertices are called *states* and whose directed edges are called *transitions*
- A start state $s \in Q$
- A set of accepting states $A \subseteq Q$
- An *input* function $\iota: E \rightarrow \Sigma \cup \{\varepsilon\}$
- An *output* function $\omega: E \rightarrow \Lambda \cup \{\varepsilon\}$

An *accepting walk* in an NFT M is an alternating sequence of vertices and directed edges

$$W = (q_0, e_1, q_1, e_2, q_2, \dots, e_n, q_n)$$

where $q_0 = s$ and $q_n \in A$, and for each index i , edge e_i goes from state q_{i-1} to state q_i . Every accepting walk has an *input label*

$$\iota(W) = \iota(e_1) \cdot \iota(e_2) \cdot \dots \cdot \iota(e_n) \in \Sigma^*$$

and an *output label*

$$\omega(W) = \omega(e_1) \cdot \omega(e_2) \cdot \dots \cdot \omega(e_n) \in \Lambda^*.$$

Finally, for any strings $w \in \Sigma^*$ and $\omega \in \Lambda^*$, we say that M *transforms* w into ω , written $w \Rightarrow_M \omega$, if there is a single accepting walk W such that $\iota(W) = w$ and $\omega(W) = \omega$.

Let M be an arbitrary NFT, and let $L \subseteq \Sigma^*$ and $L' \subseteq \Lambda^*$ be arbitrary regular languages. Prove that the language

$$\begin{aligned} M(L, L') &= \{w \in L \mid w \Rightarrow_M \omega \text{ for some } \omega \in L'\} \\ &= \{\iota(W) \mid \text{accepting walk } W \text{ in } M \text{ such that } \iota(W) \in L \text{ and } \omega(W) \in L'\} \end{aligned}$$

is also regular. (Most of the languages in problems 7 and 10 can be described this way.)

[None of the terminology or notation in this question is standard.]