

# The Complexity of Downward Closures of Indexed Languages

Richard Mandel 

Max Planck Institute for Software Systems (MPI-SWS), Kaiserslautern, Germany

Corto Mascle 

Max Planck Institute for Software Systems (MPI-SWS), Kaiserslautern, Germany

Georg Zetsche 

Max Planck Institute for Software Systems (MPI-SWS), Kaiserslautern, Germany

---

## Abstract

Indexed languages are a classical notion in formal language theory, which has attracted attention in recent decades due to its role in higher-order model checking: They are precisely the languages accepted by order-2 pushdown automata.

The downward closure of an indexed language – the set of all (scattered) subwords of its members – is well-known to be a regular over-approximation. It is known since 2015 that the downward closure of a given indexed language is effectively computable. However, the algorithm comes with no complexity bounds, and it has remained open whether a primitive-recursive construction exists.

We settle this question and provide a triply (resp. quadruply) exponential construction of a non-deterministic (resp. deterministic) automaton. We also prove (asymptotically) matching lower bounds. For the upper bounds, we rely on recent advances in semigroup theory, which let us compute bounded-size summaries of words with respect to a finite semigroup. By replacing stacks with their summaries, we are able to transform an indexed grammar into a context-free one with the same downward closure, and then apply existing bounds for context-free grammars.

**2012 ACM Subject Classification** Theory of computation → Regular languages

**Keywords and phrases** Higher-order pushdown automata, well quasi-orders, semigroup algebra

**Digital Object Identifier** 10.4230/LIPIcs.LICS.2026.69

**Related Version** *Full Version*: <https://arxiv.org/abs/2601.19466> [47]

**Funding** Funded by the European Union (ERC, FINABIS, 101077902). Views and opinions expressed are however those of the author(s) only and do not necessarily reflect those of the European Union or the European Research Council Executive Agency. Neither the European Union nor the granting authority can be held responsible for them.



**Acknowledgements** We are grateful to Alexander Kartzow for answering questions about his work [40], and to Hermann Gruber for making us aware of literature on pumping constants.

## 1 Introduction

**Downward closures.** For finite words  $u$  and  $v$ , we say that  $u$  is a (*scattered*) *subword* of  $v$ , written  $u \preceq v$ , if  $v$  can be obtained from  $u$  by inserting letters. The *downward closure* of a language  $L \subseteq \Sigma^*$  is the set  $L\downarrow = \{u \in \Sigma^* \mid u \preceq v \text{ for some } v \in L\}$  of all subwords of members of  $L$ . By the well-known Higman’s lemma [35], the downward closure  $L\downarrow$  is regular for *any* language  $L$ .

This makes the downward closure a fundamental abstraction – it is a regular overapproximation that preserves information about pattern occurrences and unboundedness behavior. Specifically, in the verification of complex systems, downward closures are often used to replace infinite-state components by finite-state ones, which then enables the algorithmic



© Richard Mandel, Corto Mascle, and Georg Zetsche;  
licensed under Creative Commons License CC-BY 4.0

41st Annual Symposium on Logic in Computer Science (LICS 2026).

Editors: Claudia Faggian and Joost-Pieter Katoen; Article No. 69; pp. 69:1–69:28



Leibniz International Proceedings in Informatics

Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany

analysis of the entire system. Examples include asynchronous programs [16, 46], shared-memory systems with dynamic thread creation [3, 8, 15, 17, 18], parameterized asynchronous shared-memory systems [44], and systems communicating via lossy channels [1, 4, 9].

For these reasons, it is often useful to *compute downward closures*, which means constructing a finite automaton for  $L\downarrow$  when given a description (i.e. a grammar or an infinite-state recognizer) for  $L$ . This is a notoriously difficult task, as it requires a deep understanding of how  $L$  is generated/recognized. Therefore, the problem of computing downward closures has attracted significant attention over recent decades, with work on context-free languages [22, 57], systems with counting and concurrency [3, 10, 11, 31, 58], models of higher-order recursion [20, 32, 59], lossy channel systems [1, 48], general algorithms for broad classes of infinite-state systems [6, 59], representation sizes [12, 18, 30, 46], related algorithmic tasks [5, 25–27, 50, 51, 60], and even computability beyond subwords [6, 7, 13, 23, 61].

**Indexed languages.** A setting where downward closure computation is particularly challenging is that of *indexed languages* [2], a classical notion in formal language theory that generalizes context-free languages. Essentially, indexed grammars differ from context-free grammars in that each non-terminal carries a stack, which can be pushed and popped through special rules. These grammars have recently attracted interest because of their role in higher-order model-checking [43, 49]: Indexed grammars are equivalent to order-2 pushdown automata, from the hierarchy of *higher-order pushdown automata (HOPA)*, which model safe higher-order recursion [41, 42] (see also the survey [49]). Level  $k$  of this hierarchy consists of the *order- $k$  pushdown automata ( $k$ -PDA)*, which have access to *stacks of ( $\dots$ ) stacks* – with nesting depth  $k$ : In particular, a 1-PDA is an ordinary pushdown automaton, whereas a 2-PDA has a stack of stacks, where it can operate on the top-most stack as an ordinary pushdown automaton; but it can also copy the top-most stack.

There is a significant gap in our understanding of downward closures of indexed languages (and HOPA more generally). There is a general approach for computing downward closures [59]. Based on this approach, computability of downward closures has been shown for indexed languages (equivalently, 2-PDA) [59], then for general HOPA [32], and even higher-order recursion schemes [20].

However, while downward closure computability is settled for these models, the complexity has remained a long-standing open problem. This is because the algorithm from [59] enumerates automata and then solves instances of the so-called *simultaneous unboundedness problem (SUP)* to decide whether the current automaton in fact recognizes  $L\downarrow$ . Thus, although the complexity of the SUP itself is well-understood for HOPA [14, 50], there are no complexity upper bounds for the entire task of computing downward closures of indexed languages (nor for languages recognized by HOPA in general). In fact, even the existence of a primitive-recursive (or even hyper-Ackermannian) upper bound remained open<sup>1</sup>.

**Contribution.** In this work, we settle the complexity of computing downward closures of indexed languages. We show that given an indexed language  $L$ , one can compute a non-deterministic finite automaton (NFA) for  $L\downarrow$  in triply exponential time (and hence of triply exponential size). We also provide a triply exponential lower bound, improving on the doubly exponential lower bound in [60].

---

<sup>1</sup> Given that computing downward closures involves a well-quasi-ordering (WQO) at its core, and WQO-based algorithms employing the subword ordering can often be furnished with hyper-Ackermannian/multiply-recursive upper bounds via length-function theorems [53, 54], one might hope for such a bound here. Unfortunately, it is not clear how to apply length-function theorems to downward closure computation.

Furthermore, our constructions also provide a tight bound for the computation of a deterministic finite automaton (DFA) for  $L_\downarrow$ : It follows that one can construct a quadruply-exponential-sized DFA for  $L_\downarrow$ , and we provide a quadruply-exponential lower-bound.

Moreover, our results settle the complexity of decision problems involving downward closures of indexed languages. The *downward closure inclusion problem* asks, whether for given indexed languages  $L_1$  and  $L_2$ , we have  $L_1\downarrow \subseteq L_2\downarrow$ . Similarly, the *downward closure equivalence problem* asks whether  $L_1\downarrow = L_2\downarrow$ . We show that both problems are co-3-NEXP-complete.

Finally, our construction for the downward closure lower bound also settles another question about indexed languages: It implies that the known triply exponential upper bound on the pumping threshold for indexed grammars [34, 56] (the maximal length of words in a finite indexed language) is in fact asymptotically tight. (In [40], a doubly exponential upper bound was claimed. However, as confirmed by the author of [40], that is a miscalculation; see Section 3 and the full version [47].)

**Why are the results unexpected?** The complexity results come as a considerable surprise. This is because (tight) complexity bounds for HOPA are usually towers of exponentials where the height grows linearly with the order. For example, for  $k \geq 1$ , the emptiness problem for  $k$ -PDA is  $(k - 1)$ -EXPTIME-complete [24, comment before Thm. 7.12] (note that 0-EXPTIME = P). The same is true of the SUP [50, Thm. 3]. Since downward closure inclusion and equivalence are coNP-complete for NFAs (see [12, Sec. 5] and [39, Prop. 7.3]) and co-1-NEXP-complete for 1-PDAs [60, Tab. 1], and hence co- $k$ -NEXP-complete for  $k$ -PDAs with  $k \leq 1$ , one would expect co-2-NEXP rather than co-3-NEXP for 2-PDAs. In fact, the best (and only) known lower bound until now has been co-2-NEXP [60, Cor. 18]. Similarly, since downward closure NFAs are polynomial-sized for given NFAs (sometimes considered 0-PDAs) and exponential-sized for given 1-PDAs [12, Cor. 6], one would expect a doubly exponential bound for 2-PDAs. In fact, the best known lower bound on the NFA size for downward closures had been doubly exponential [60, remarks before Cor. 18].

**Key ingredients.** Indexed grammars extend context-free grammars by equipping the nodes of derivation trees with pushdown store (which we also call stack). This way, each branch of a derivation tree corresponds to a run of a pushdown automaton.

Our construction relies on recent advances in finite semigroup theory, which provide succinct “summaries” of arbitrarily long words relative to a finite semigroup [28]. We use this (for a suitable semigroup) to replace the stack in each derivation tree node by such a summary, each of which uses exponentially many bits. This replacement changes the overall language, but preserves the downward closure. Once the information in each node is bounded, we can transform the grammar into a doubly-exponential-sized context-free grammar. This yields the triply exponential upper bound overall, since for context-free grammars, existing algorithms yield exponential-sized downward closure NFAs [12, Corollary 6].

The aforementioned summaries are similar in spirit to Simon’s factorization forests [55]. The latter annotate words by trees: Relative to a morphism  $\varphi: \Sigma^* \rightarrow M$  into a finite monoid  $M$ , a factorization forest for  $w \in \Sigma^*$  is a tree of height bounded by a function of  $|M|$  that allows evaluating  $\varphi$  on infixes of  $w$  without processing the entire infix. Here, a crucial idea is that a sequence of infixes that all map under  $\varphi$  to the same idempotent  $e$  must evaluate to  $e$ .

Similar to factorization forests, our summaries also exploit repetitions of idempotents to collapse long infixes in a stack. However, in contrast to factorization forests, our summaries reduce the stack word to one of bounded length (hence losing information). Also crucially, taking summaries is compatible with pushing stack symbols: Given a summary of  $w$ , one can compute a summary of  $aw$ .

**Structure of the paper.** We recall necessary notations and basic results in Section 2, and state the main results in Section 3. Then in Section 4 we show how to make a given indexed grammar *productive*, meaning that every partial computation can be extended to a full one. In Section 5 we introduce an object at the core of our construction, the *production monoid*. In Section 6 we introduce new rules that extend and reduce parts of the stack without altering the downward closure of the language. We then proceed in Section 7 to define *summaries* of stack contents with respect to the production monoid. In Section 8 we leverage those summaries to compute from an indexed grammar a context-free one with the same downward closure, to which known constructions apply. Finally, in Section 9 we complement our upper bounds with matching lower bounds.

*This paper contains internal links; every occurrence of a term is linked to its definition. The reader can click on terms (and some notations), or simply hover over them on some pdf readers, to get their definition.*

*Due to space constraints, most of the proofs are deferred to the full version [47].*

## 2 Preliminaries

**Words, trees and languages.** Given a finite alphabet  $\Sigma$ , we write  $\Sigma^*$  for the set of words over  $\Sigma$ , and  $\Sigma^k$  for the set of words of length  $k$ . Given  $w \in \Sigma^*$ , we denote  $|w|$  its length. We assume familiarity with basic finite automata theory, see [52] for an introduction.

We write  $u \preceq v$  if  $u$  is a (scattered) *subword* of  $v$ ; that is, if  $u$  can be obtained from  $v$  by removing some of its letters. If a word  $w$  is equal to  $uv$ , then we call  $u$  a *prefix*, and  $v$  a *suffix* of  $w$ . The *downward closure* of a language  $L \subseteq \Sigma^*$  is the set of subwords of words of  $L$ , and is denoted  $L\downarrow$ . An important reason for studying downward closures is that they are always regular. This was first shown by Haines [33, Theorem 3], but also follows from  $\preceq$  being a well-quasi ordering, which was shown earlier by Higman [35, Theorem 4.3]:

► **Theorem 2.1.** *The downward closure of any language is regular.*

A finite ordered  $\Sigma$ -labeled binary tree is a pair  $(\tau, \lambda)$  with  $\tau$  a finite prefix-closed subset of  $\{0, 1\}^*$  such that for all  $\nu \in \tau$ ,  $\nu 1 \in \tau$  implies  $\nu 0 \in \tau$ , and  $\lambda : \tau \rightarrow \Sigma$  a function mapping each node in  $\tau$  to a label. We use the usual terminology for trees, with *node*, *leaf*, *branch*, *subtree*, *parent*, *child*, *ancestor*, *descendant* etc. retaining their usual meaning. The *leaf word* of  $\tau$  is the word obtained by concatenating  $\lambda(\nu_1) \dots \lambda(\nu_k)$ , where  $\nu_1, \dots, \nu_k$  are the leaves of  $\tau$  read from left to right, i.e., in lexicographic order.

**Indexed grammars.** To simplify definitions and proofs, we use a syntax for indexed grammars that resembles the Chomsky normal form used for context-free grammars.

An *indexed grammar* is a tuple  $\mathcal{G} = (N, T, I, P, S)$  with

- $N$  the set of non-terminals, and  $S \in N$  the starting non-terminal
- $T$  the set of terminal symbols
- $I$  the set of index symbols, also called stack symbols
- $P$  a set of productions, which are of the following types:
  - $A \rightarrow w$  with  $w \in T^*$ .
  - $A \rightarrow BC$  with  $A, B, C \in N$
  - $A \rightarrow Bf$  with  $A, B \in N$  and  $f \in I$
  - $Af \rightarrow B$  with  $A, B \in N$  and  $f \in I$

We define the *size* of  $\mathcal{G}$  as  $|N| + |P| + \sum_{A \rightarrow w \in P} |w|$ . A *context-free grammar* (or CFG) is an indexed grammar where all productions are of the two first forms.

A *sentential form* is a word over the (infinite) alphabet  $NI^* \cup T$ . The set of sentential forms is denoted  $\mathbf{SF}_{\mathcal{G}}$ , or simply  $\mathbf{SF}$  when the grammar is clear from context. We write the elements of  $NI^*$  as  $A[z]$ , with  $A \in N$  a non-terminal and  $z \in I^*$  interpreted as the content of its stack; such an element is sometimes referred to as a *term*. If  $u \in \mathbf{SF}$ , we occasionally use the notation  $u[z]$  to denote the sentential form obtained by pushing  $z$  on top of the stack of every non-terminal in  $u$ . The derivation relation  $\Rightarrow_{\mathcal{G}}$  (or simply  $\Rightarrow$ ) over  $\mathbf{SF}$  is defined as:

$$\begin{aligned} uA[z]v &\Rightarrow uB[z]C[z]v && \text{if } A \rightarrow BC \in P, \\ uA[z]v &\Rightarrow uB[fz]v && \text{if } A \rightarrow Bf \in P, \\ uA[fz]v &\Rightarrow uB[z]v && \text{if } Af \rightarrow B \in P, \\ uA[z]v &\Rightarrow uvv && \text{if } A \rightarrow w \in P. \end{aligned}$$

for all  $u, v \in \mathbf{SF}$ ,  $A, B, C \in N$ ,  $f \in I$ ,  $z \in I^*$  and  $w \in T^*$ . We write  $\Rightarrow_{\mathcal{G}}^*$  (or simply  $\Rightarrow^*$ ) for the reflexive transitive closure of  $\Rightarrow_{\mathcal{G}}$ .

The subword relation  $\preceq$  is extended to  $\mathbf{SF}$  in the natural way, i.e.  $u \preceq v$  if  $u$  can be obtained from  $v$  by deleting terms. Note that on  $\mathbf{SF}$ , the ordering  $\preceq$  is not a well-quasi ordering, since any two terms  $A[z]$  and  $A[z']$  are incomparable for  $z, z' \in I^*$ ,  $z \neq z'$ .

Given a sentential form  $u$ , we write  $\mathbf{L}_{\mathbf{SF}}(u)$  for the set of sentential forms derivable from  $u$ ,  $\{v \in \mathbf{SF} \mid u \Rightarrow^* v\}$ . The *language* of  $\mathcal{G}$  is denoted  $\mathbf{L}(\mathcal{G})$  and defined as  $\mathbf{L}_{\mathbf{SF}}(S) \cap T^*$ . Furthermore, for all  $X \subseteq N$  and  $u \in \mathbf{SF}$ , the language  $\mathbf{L}_X(u)$  is defined as the set  $\mathbf{L}_{\mathbf{SF}}(u) \cap (X \cup T)^*$  of sentential forms derivable from  $u$  with all stacks empty, and all non-terminals in  $X$ . In particular,  $\mathbf{L}_{\emptyset}(u)$  is the set of terminal words which can be derived from  $u$ . If the language  $\mathbf{L}_{\emptyset}(u)$  is non-empty then we say that  $u$  is *productive*.

A *derivation tree* is a finite ordered tree  $\tau$  whose nodes are labeled by elements of  $NI^* \cup T^*$ , with the following constraints. For each node  $\nu$  with label  $A[z]$  (where  $A \in N$  and  $z \in I^*$ ), exactly one of the following holds

- $\nu$  is a leaf
- $\nu$  has one child labeled  $w$  for some  $A \rightarrow w \in P$
- $\nu$  has two children labeled  $B[z]$  and  $C[z]$ , for some production  $A \rightarrow BC \in P$
- $\nu$  has one child labeled  $B[fz]$  for some  $A \rightarrow Bf \in P$
- $\nu$  has one child labeled  $B[z']$  for some  $Af \rightarrow B \in P$ , with  $z = fz'$ .

A derivation tree whose root is labeled  $A[z]$  and whose leaf word is equal to  $u$  is called a *derivation tree from  $A[z]$  to  $u$* . If  $u \in T^*$  we say that the derivation tree is *complete*.

► **Remark 2.2.** An easy induction shows that for all  $A[z] \in NI^*$ , a sentential form  $u$  can be derived from  $A[z]$  if and only if  $u$  is the leaf word of a derivation tree whose root is labeled  $A[z]$ . In the forthcoming proofs we will either use sentential forms or derivation trees depending on what is most convenient.

► **Remark 2.3.** When describing examples of indexed grammars, we will use rules of the form  $Af \rightarrow u$  and  $A \rightarrow u$  with  $u \in (N \cup T)^*$ . This is just syntactic sugar, as we can replace them with rules from the definition of indexed grammar while adding a small number of non-terminals, in the spirit of the Chomsky normal form [19]. This transformation incurs a linear size increase of the grammar. Details are left to the full version [47].

► **Example 2.4.** We can define the language  $\{a^n b^{n^2} \mid n \in \mathbb{N}\}$  with an indexed grammar:

$$\begin{array}{ll}
 S \rightarrow Tg & \# g \text{ is the stack bottom symbol;} \\
 T \rightarrow Tf & \# \text{ we push some number } n \text{ of } f; \\
 T \rightarrow A & \\
 Ag \rightarrow \varepsilon & \# \text{ if } n = 0 \text{ then return } \varepsilon; \\
 Af \rightarrow C & \# \text{ if } n > 0 \text{ we pop an } f; \\
 C \rightarrow aAB & \# \text{ repeat } n \text{ times to get } a^n B[g]B[fg] \cdots B[f^{n-1}g]; \\
 Bf \rightarrow bbB & \\
 Bg \rightarrow b & \# \text{ the } Bs \text{ output } b^{\sum_{i=0}^{n-1} (2i+1)} = b^{n^2}.
 \end{array}$$

► **Remark 2.5.** We can assume without loss of generality that every symbol pushed on the stack carries the information of the production rule used to push it (this property can always be ensured by adding a quantity of new stack symbols and production rules that is at most quadratic in the size of the grammar). Formally, we assume that there are functions  $\alpha, \beta : I \rightarrow N$  such that for every push rule  $A \rightarrow Bf$  we have  $A = \alpha(f)$  and  $B = \beta(f)$ . We also assume that for all  $f \in I$  there is a rule of the form  $A \rightarrow Bf$  in  $P$ . Clearly stack symbols not satisfying this condition can be removed.

**Complexity.** We define the functions  $\exp_k : \mathbb{N} \rightarrow \mathbb{N}$  inductively by setting  $\exp_0(n) = n$  and  $\exp_{k+1}(n) = 2^{\exp_k(n)}$ . A function  $f : \mathbb{N} \rightarrow \mathbb{N}$  is (at most)  $k$ -fold exponential if there is a constant  $c > 0$  such that  $f(n) \leq \exp_k(n^c)$  for almost all  $n$ . We say that  $f$  is at least  $k$ -fold exponential if there is a constant  $c > 0$  such that  $f(n) \geq \exp_k(n^c)$  for infinitely many  $n$ . Instead of 2-fold, 3-fold, 4-fold exponential, resp., we also say doubly, triply, or quadruply exponential, resp. For  $k \geq 1$ , the class **co- $k$ -NEXP** consists of the complements of sets accepted by an  $f$ -time-bounded non-deterministic Turing machine, for some  $f : \mathbb{N} \rightarrow \mathbb{N}$  that is at most  $k$ -fold exponential.

### 3 Main results

In this section, we present the main results of this work.

**Non-deterministic automata.** Our first main contribution is an algorithm to compute an NFA of at most triply exponential size for the downward closure of an indexed language:

► **Theorem 3.1.** *Given an indexed grammar  $\mathcal{G}$ , one can compute (in triply exponential time) a triply-exponential-sized NFA for  $L(\mathcal{G})_\downarrow$ .*

We also provide an asymptotically matching lower bound, which we infer from the following result:

► **Theorem 3.2.** *Given  $n \in \mathbb{N}$  (in unary encoding), we can compute in polynomial time an indexed grammar for the language  $\{a^{\exp_3(n)}\}$ .*

An NFA for  $\{a^{\exp_3(n)}\}_\downarrow$  clearly requires at least  $\exp_3(n)$  states, implying Corollary 3.3.

► **Corollary 3.3.** *There is a family  $(\mathcal{G}_n)_{n \geq 1}$  of indexed grammars of size polynomial in  $n$  such that any NFA for  $L(\mathcal{G}_n)_\downarrow$  requires at least  $\exp_3(n)$  states.*



**Deterministic automata.** Of course, Theorem 3.1 implies a quadruply exponential upper bound for deterministic automata:

► **Corollary 3.4.** *Given an indexed grammar  $\mathcal{G}$ , one can compute (in quadruply exponential time) a DFA of at most quadruply exponential size for  $L(\mathcal{G})\downarrow$ .*

Here, we have an asymptotically matching lower bound as well:

► **Theorem 3.5.** *There is a family  $(\mathcal{G}_n)_{n \geq 1}$  of indexed grammars of size polynomial in  $n$  such that any DFA for  $L(\mathcal{G}_n)\downarrow$  requires at least  $\exp_4(n)$  states.*

**Downward closure comparisons.** Theorem 3.1 and our construction for Theorem 3.2 also allow us to settle the complexity of decision problems related to downward closures. The *downward closure inclusion problem* (for indexed languages) asks whether two given indexed languages  $L_1, L_2$  satisfy  $L_1\downarrow \subseteq L_2\downarrow$ . Similarly, the *downward closure equivalence problem* asks whether  $L_1\downarrow = L_2\downarrow$ .

In [60, Corollary 18], it was shown that downward closure inclusion and equivalence for indexed languages are co-2-NEXP-hard. Here, we settle their precise complexity:

► **Theorem 3.6.** *Downward closure inclusion and downward closure equivalence for indexed languages are co-3-NEXP-complete.*

Here, the upper bounds follow from Theorem 3.1 and the fact that downward closure inclusion and equivalence are coNP-complete for NFAs (see [12, Section 5] and [39, Proposition 7.3]).

**The pumping threshold for indexed languages.** Our lower bound technique also settles the growth of the pumping threshold of indexed grammars. Consider the function

$$\mathfrak{P}(n) = \max\{|w| \mid \text{there is an indexed grammar } \mathcal{G} \text{ of size } \leq n \\ \text{with } w \in L(\mathcal{G}) \text{ such that } L(\mathcal{G}) \text{ is finite}\}$$

We call  $\mathfrak{P}(n)$  the *pumping threshold* (for size  $n$ ) because placing an upper bound on  $\mathfrak{P}(n)$  usually involves a pumping argument.

An analogous pumping threshold function for NFAs and CFGs is well-understood: It is linear for NFAs (see [36] for related results) and exponential for CFGs (exact bounds can be found in [29]). For indexed grammars, there are two proofs of a triply exponential upper bound: the pumping lemmas of Hayashi [34, Theorem 5.1] and Smith [56, Theorem 1] (Smith's proof mentions the bound explicitly; Hayashi's proof requires some analysis for this). We complement this by showing a triply exponential lower bound:

► **Corollary 3.7.**  *$\mathfrak{P}$  grows at least triply exponentially.*

This follows from Theorem 3.2, because  $\{a^{\exp_3(n)}\}$  is finite but has an indexed grammar of size polynomial in  $n$ . It should be noted that [40, Section 7] claims a doubly exponential upper bound for  $\mathfrak{P}$ . Unfortunately, as confirmed by the author of [40], that is due to a miscalculation, see the full version [47] for a discussion.

► **Remark 3.8.** A triply exponential upper bound on  $\mathfrak{P}(n)$  also follows from our Theorem 3.1: If an indexed grammar  $\mathcal{G}$  generates a word that is longer than the number of states of an NFA for  $L(\mathcal{G})\downarrow$ , then  $L(\mathcal{G})\downarrow$  must be infinite, and hence also  $L(\mathcal{G})$ .

**Structure of the paper.** In Sections 4–8, we will prove Theorem 3.1, from which all remaining upper bounds follow. In Section 9, we will then prove all lower bounds.

#### 4 Productiveness

For the rest of the paper, we assume that our input grammar  $\mathcal{G}$  generates a non-empty language. This is because emptiness of indexed languages can be decided in EXPTIME [24, Theorem 7.12], and if  $L(\mathcal{G})$  is empty, an NFA for  $L(\mathcal{G})\downarrow$  is immediate.

However, we will need to establish the stronger guarantee of *productiveness*, which expresses the absence of deadlocks. This means that if we can produce a sentential form  $u$ , then from  $u$  we can derive a terminal word in  $T^*$ .

**Productive grammars.** We say that  $\mathcal{G}$  is *productive* if for all  $u \in L_{SF}(S)$  we have  $L_\emptyset(u) \neq \emptyset$ , that is, from every sentential form obtained from  $S$  we can derive a word in  $T^*$ .

This property is especially useful for downward closure computation. In a productive indexed grammar, we can observe: for all  $u \in L_{SF}(S)$ , for all  $v$  such that  $v \preceq u$ , we have  $L_\emptyset(v) \subseteq L_\emptyset(u)\downarrow$ . In other words, we can interleave derivations and deletion of terms without any risk of obtaining “extra” terminal words. This property does not hold in general, because deleting terms that cannot produce terminal words may allow the derivation of a terminal word that is not in the downward closure.

► **Example 4.1.** The following grammar  $\mathcal{G}$  is not productive.

$$\begin{array}{lll} S \rightarrow S\perp & S \rightarrow Sf & S \rightarrow Sg \\ S \rightarrow AA & S \rightarrow AAC & \\ Af \rightarrow aA & Ag \rightarrow b & \\ Cf \rightarrow cC & Cg \rightarrow CAB & \\ A\perp \rightarrow \varepsilon & C\perp \rightarrow \varepsilon & \end{array}$$

The language of  $\mathcal{G}$  is  $\{w^2 \mid w \in \{a, b\}^*\} \cup \{a^{2n}c^n \mid n > 0\}$ . In particular, by setting  $u = A[gf\perp]A[gf\perp]C[f\perp]A[f\perp]B[f\perp]$  and  $v = C[f\perp]A[f\perp]$  we have  $v \preceq u$  and  $u \in L_{SF}(S)$ , but  $ca \in L_\emptyset(v)$  while  $L_\emptyset(u)\downarrow = \emptyset$  (since  $B$  cannot produce a terminal word). Moreover,  $ca \notin L(\mathcal{G})\downarrow$ . In this case, it is easy to modify  $\mathcal{G}$  to obtain a productive grammar which generates the same language.

**Tracking productivity of non-terminals.** Before we describe our method for achieving productiveness, we observe that tracking the non-terminals which, with a given stack content, can derive a sentential form with terms confined to a certain subset, gives rise to a left semigroup action.

► **Definition 4.2.** For  $X \subseteq N$  and  $z \in I^*$ , we define

$$z \cdot X = \{A \in N \mid L_X(A[z]) \neq \emptyset\},$$

i.e.  $z \cdot X$  is the set of non-terminals  $A$  so that the term  $A[z]$  can derive a sentential form consisting of terminals and empty-stack occurrences of non-terminals in  $X$ . Let  $U$  be defined as the set of non-terminals which can derive a word in  $T^*$ , i.e.

$$U = \{A \in N \mid L_\emptyset(A) \neq \emptyset\}.$$

Note that  $L(\mathcal{G})$  is empty if and only if  $S \notin U$ .

We will often rely on the fact that  $I^*$  acts (on the left) as a semigroup<sup>2</sup> on the power set  $2^N$ , as we state now:

<sup>2</sup> However, it does not act as a monoid, because  $\varepsilon \cdot X$  is not necessarily  $X$ , as every set  $z \cdot X$  includes  $U$ .



► **Lemma 4.3.** *For every  $f \in I$ ,  $z \in I^*$ , and  $X \subseteq N$ , we have  $fz \cdot X = f \cdot (z \cdot X)$ . Moreover,  $z \cdot \mathcal{U} = z \cdot \emptyset$ .*

Here, only the inclusion  $fz \cdot X \subseteq f \cdot (z \cdot X)$  is not trivial: It holds because a derivation that eliminates  $fz$  from the stack must in each branch first remove  $f$ , and then  $z$ .

**Achieving productiveness.** We are now ready to present our construction of a productive equivalent of  $\mathcal{G}$ . Formally, the *annotated version* of  $\mathcal{G}$ , denoted  $\bar{\mathcal{G}} = (\bar{N}, T, \bar{I}, \bar{P}, \bar{S})$ , is defined as follows. The general idea is to integrate in to label every point in the stack with a set of non-terminals, which are the non-terminals  $A$  such that if we take the stack content  $z'$  from that point, one can derive a terminal word from  $A[z']$ . By maintaining this information, we avoid producing terms that cannot derive a terminal word.

Recall that we assumed  $L(\mathcal{G}) \neq \emptyset$ , thus  $S \in \mathcal{U}$ . Otherwise, we have:

- $\bar{N} = \{(A, X) \mid X \subseteq N, A \in X\}$ ,  $\bar{I} = I \times 2^N$ , and  $\bar{S} = (S, \mathcal{U})$ ,
- $\bar{P}$  contains the following rules:
  - $(A, X) \rightarrow w$  for all  $A \rightarrow w \in P$  with  $A \in X$  and  $w \in T^*$
  - $(A, X) \rightarrow (B, X)(C, X)$  for all  $A \rightarrow BC \in P$  with  $A, B, C \in X$
  - $(A, X) \rightarrow (B, Y)(f, X)$  for all  $A \rightarrow Bf \in P$  with  $Y = f \cdot X$ ,  $A \in X$  and  $B \in Y$
  - $(A, Y)(f, X) \rightarrow (B, X)$  for all  $Af \rightarrow B \in P$  with  $Y = f \cdot X$ ,  $A \in Y$  and  $B \in X$

Note that a stack word  $\bar{z} = (f_n, X_n) \cdots (f_1, X_1) \in \bar{I}^*$  appearing as an infix of a stack content in a derivation of  $\bar{\mathcal{G}}$  must be so that  $X_i = f_{i-1} \cdot X_{i-1}$  for all  $i > 1$ . A stack word satisfying this property is determined by its projection onto  $I^*$  and its last element. Given  $z = f_n \cdots f_1 \in I^*$  and  $X \subseteq N$ , define  $\bar{z}^X$  as  $(f_n, X_n) \cdots (f_1, X_1)$  with  $X_1 = X$  and  $X_{i+1} = f_i \cdot X_i$  for all  $i < n$ . Given  $A[z] \in NI^+$  with  $z = f_n \cdots f_1$  and  $A \in z \cdot X$ , the *X-based annotation* of  $A[z]$  is the  $\bar{\mathcal{G}}$  term  $(A, Y)[\bar{z}] = (A, Y)[(f_n, X_n) \cdots (f_1, X_1)] \in \bar{N} \bar{I}^+$  such that  $X_1 = X$ ,  $Y = f_n \cdot X_n$  and  $X_i = f_{i-1} \cdot X_{i-1}$  for all  $i > 1$ . In particular, by Lemma 4.3 we have  $X_i = f_{i-1} \cdots f_1 \cdot X$  for all  $i$  and  $Y = z \cdot X$ . In the case of an empty stack, the *X-based annotation* of  $A$  is  $(A, X)$ .

**Correctness of the construction.** Having defined  $\bar{\mathcal{G}}$ , we can prove that it serves its purpose:

► **Lemma 4.4.**  *$\mathcal{G}$  and  $\bar{\mathcal{G}}$  have the same language, and  $\bar{\mathcal{G}}$  is productive.*

► **Remark 4.5.** Although we are able to turn any indexed grammar into a productive one with the same language, this transformation incurs an exponential blow-up in the size. Therefore, in order to obtain tight complexity bounds, we do not simply assume productiveness of the input grammar. Rather, we work with annotated versions explicitly when needed.

## 5 The stack monoid

**Overall goal.** Another key aspect of our construction is the *stack monoid* – a finite monoid in which we evaluate stack contents. Essentially, we map a stack  $\bar{z} \in \bar{I}^*$  to a monoid element that encodes<sup>3</sup>

- the non-terminals from and to which this content was pushed (which are unique thanks to Remark 2.5), and
- for each pair  $A, B$  of non-terminals, whether we can derive a *sentential form* containing  $B$  from  $A[\bar{z}]$  (i.e., whether we can obtain a  $B$  by popping  $\bar{z}$  from  $A$ ).

<sup>3</sup> Notice that there is a dissymmetry between push and pop here (in fact, throughout the construction). This is because a stack symbol is only pushed once, but may be popped on multiple branches.

The purpose of this encoding is that based on this information, we will be able to simplify (or expand) stack contents during derivations, *while preserving the downward closure*. For example, if we have a derivation from  $A[z]$  to a **sentential form**  $uBv$ , then we could just erase  $u$  and  $v$ , thus turning  $A[z]$  into  $B$ . Even better, if we have a derivation from  $A[z]$  to  $A$  then, roughly speaking, this means we can turn any term  $A[zz']$  into  $A[z']$ .

Observations like this, based on the stack monoid, will be used in Section 6 to define more involved stack manipulations. In Section 7, these will allow us to reduce, roughly speaking, every stack to one of doubly exponentially many.

**Semigroup terminology.** We begin with some terminology. A semigroup  $(\mathbf{S}, \cdot)$  is a set equipped with an associative product. We will often identify a semigroup with its set of elements and denote it simply as  $\mathbf{S}$ , when the product operation is clear. A monoid  $(\mathbf{M}, \cdot, \mathbf{1}_\mathbf{M})$  is a semigroup with a neutral element  $\mathbf{1}_\mathbf{M}$ .

For each monoid  $(\mathbf{M}, \cdot)$ , we define  $\psi_\mathbf{M} : \mathbf{M}^* \rightarrow \mathbf{M}$  to be its *evaluation morphism*, which maps a sequence of elements of  $\mathbf{M}$  to its product, with the empty sequence being mapped to  $\mathbf{1}_\mathbf{M}$ . An element  $x \in \mathbf{M}$  is *idempotent* if  $x \cdot x = x$ . The set of idempotent elements of  $\mathbf{M}$  is denoted  $\text{Idem}(\mathbf{M})$ .

In all that follows we fix an indexed grammar  $\mathcal{G} = (N, T, I, P, S)$ . We will mainly work with its annotated version  $\bar{\mathcal{G}}$ . Note that we cannot use the annotation construction as a black box and simply assume that the input grammar is productive (see Remark 4.5).

**Formal definition.** Let us now start by describing formally the monoid at hand. It is a bit more involved than what is described above since we have to account for the productiveness issues explained in the previous section. We will use the boolean matrix monoid over non-terminals of  $\mathcal{G}$ : This monoid is defined as the set of matrices  $\mathbb{B}^{N \times N}$ , with  $\mathbb{B} = \{\top, \perp\}$ . The product is the usual product of matrices over the Boolean semiring  $(\mathbb{B}, \vee, \wedge)$ , and the neutral element is the matrix with  $\top$  on the diagonal and  $\perp$  everywhere else. Given matrices  $M_1, M_2$ , we write  $M_1 M_2$  for their product.

We also define, for all  $X \subseteq N$ , the *reachability relation*  $R_X$  over  $N$ . It relates two non-terminals if from the first one we can produce a sentential form containing the second one, with empty stacks. Formally,

$$A R_X B \text{ if and only if there is a derivation } (A, X) \xrightarrow{*}_{\bar{\mathcal{G}}} u(B, X)v \text{ with } u, v \in \mathbf{SF}$$

(note that due to the productiveness property, requiring stack emptiness is only important for  $(B, X)$ ).

Let  $Q$  be the set of tuples  $(B, Y, M, A, X)$  with  $A, B \in N$  non-terminals,  $X, Y \subseteq N$  sets of non-terminals, and  $M \in \mathbb{B}^{N \times N}$ . Define the *stack monoid* as the monoid  $(\mathbb{M}, \cdot, \mathbf{1}_\mathbb{M})$  whose elements are  $Q \cup \{\mathbf{1}_\mathbb{M}\} \cup \{\mathbf{0}_\mathbb{M}\}$ , where  $\mathbf{0}_\mathbb{M}$  satisfies  $\mathbf{0}_\mathbb{M} \cdot x = x \cdot \mathbf{0}_\mathbb{M} = \mathbf{0}_\mathbb{M}$  for all  $x \in \mathbb{M}$ , and whose product operation is defined as follows:

$$(B_2, Y_2, M_2, A_2, X_2) \cdot (B_1, Y_1, M_1, A_1, X_1) = \begin{cases} (B_2, Y_2, M_1 M_2, A_1, X_1) & \text{if } X_2 = Y_1 \\ & \text{and } B_1 R_{X_2} A_2 \\ \mathbf{0}_\mathbb{M} & \text{otherwise.} \end{cases}$$

Let  $\alpha, \beta$  be the functions described in Remark 2.5 for  $\mathcal{G}$ . We define a morphism  $\varphi : \bar{I}^+ \rightarrow \mathbb{M}$  as follows. For each letter  $(f, X) \in I$ , set

$$\varphi(f, X) = (\beta(f), f \cdot X, M_{f, X}, \alpha(f), X),$$

where for all  $A, B \in N$ ,  $M_{f,X}(A, B) = \top$  if and only if there exist  $u, v \in (X \cup T)^*$  such that  $A[f] \xRightarrow{*}_{\mathcal{G}} uBv$ . Note that computing this matrix easily reduces to an emptiness check for an indexed grammar, which can be done in exponential time.

**Feasibility.** Let us mention some basic properties of stacks that are encoded in their image in  $\mathbb{M}$ . The first concerns whether a given stack content can be pushed: We say that a non-empty stack content (in either  $I^*$  or  $\bar{I}^*$ ) is *feasible* if, starting from some non-terminal, it can be pushed onto the stack of some non-terminal. In the case of an annotated stack content  $\bar{z} = (f_n, X_n) \cdots (f_1, X_1) \in \bar{I}^*$ , this is equivalent to the existence of a derivation

$$(\alpha(f_1), X_1) \xRightarrow{*}_{\bar{\mathcal{G}}} u(\beta(f_n), f_n \cdot X_n)[\bar{z}]v$$

with  $u, v \in \mathbf{SF}$ . The following lemma says that the stack monoid distinguishes the feasible stack contents in  $\bar{I}^*$ .

► **Lemma 5.1.** *Let  $\bar{z} = (f_n, X_n) \cdots (f_1, X_1) \in \bar{I}^*$  be a stack content. The following are equivalent:*

1.  $\bar{z}$  is feasible
2.  $\varphi(\bar{z}) \neq \mathbf{0}_{\mathbb{M}}$
3. for all  $i > 1$ ,  $X_i = f_i \cdot X_{i-1}$  and  $\beta(f_{i-1}) R_{X_i} \alpha(f_i)$ .

Note that Lemma 5.1 shows that all feasible stack contents in  $\bar{I}^*$  can be written as  $\bar{z}^X$  for some feasible  $z \in I^*$  and  $X \subseteq N$  (hence, we sometimes assume that a stack content is of this form).

Let  $(A, X) \in \bar{N}$  and  $\bar{z} = (f_n, X_n) \cdots (f_1, X_1) \in \bar{I}^*$ , we say that the term  $(A, X)[\bar{z}]$  is *feasible* if  $\bar{z}$  is feasible,  $X = f_n \cdot X_n$  and  $\beta(f_n) R_X A$ .

**Pushing between specific non-terminals.** We will now see that  $\mathbb{M}$  encodes which non-terminals allow a stack content to be pushed, and which non-terminals can result. More formally,  $\varphi(\bar{z})$  encodes all pairs of non-terminals  $(C, X), (D, Y) \in \bar{N}$  such that  $(C, X)$  can derive a sentential form  $u(D, Y)[\bar{z}]v$ .

► **Lemma 5.2.** *Let  $z \in I^+$  a non-empty stack content, let  $X \subseteq N$  and let  $(B, Y, M, A, X) = \varphi(\bar{z}^X)$ . Then for all  $C \in X$  and  $D \in Y$ , the following are equivalent:*

- $C R_X A$  and  $B R_Y D$
- there exist  $u, v \in \mathbf{SF}$  such that  $(C, X) \xRightarrow{*}_{\bar{\mathcal{G}}} u(D, Y)[\bar{z}^X]v$ .

**Popping between specific non-terminals.** Finally, our monoid even encodes popping behavior. Specifically, the matrix  $M$  inside  $\varphi(\bar{z}^X)$  tells us for which non-terminals  $D$  and  $C$ , it is possible to pop  $z$  from  $D$  to  $C$ :

► **Lemma 5.3.** *Let  $z \in I^+$  a non-empty stack content, let  $X \subseteq N$  and let  $(B, Y, M, A, X) = \varphi(\bar{z}^X)$ . The following are equivalent:*

- $M(D, C) = \top$
- $C \in X$ ,  $D \in Y$  and there exist  $u, v \in (X \cup T)^*$  such that

$$D[z] \xRightarrow{*}_{\mathcal{G}} uCv$$

- $C \in X$ ,  $D \in Y$  and there exist  $u, v \in \mathbf{SF}_{\mathcal{G}}$  such that

$$(D, Y)[\bar{z}^X] \xRightarrow{*}_{\bar{\mathcal{G}}} u(C, X)v.$$

## 6 Pumping and skipping

We introduce two new derivation rules on the terms of  $\bar{\mathcal{G}}$ , and show that they preserve the downward-closure of the resulting language. Essentially, we show that, under certain conditions, sequences of more than  $2|N|$  contiguous infixes mapping to the same idempotent can be extended and reduced.

This will let us abstract stack contents by forgetting everything but the first and last  $N$  elements in such sequences of infixes mapping to the same idempotent. By adapting a recent construction by Gimbert, Mascle and Totzke [28], we will show that this allows us to obtain summaries of stack contents, of size bounded by an exponential function in the size of  $\mathcal{G}$ .

**Pump and skip derivation steps.** Let us formally define the two rules. The *pump rule* is defined as follows:

$$(B, X)[\bar{z}] \rightarrow_{\text{pump}} (B, X)[z_e \bar{z}]$$

for all  $(A, X) \in \bar{N}$ ,  $e = (B, X, M, A, X) \in \text{Idem}(\mathbb{M}) \setminus \{\mathbf{0}_{\mathbb{M}}, \mathbf{1}_{\mathbb{M}}\}$ ,  $z_e \in \bar{I}^+$  with  $\varphi(z_e) = e$ , and  $\bar{z} \in \bar{I}^*$ .

The *skip rule* is first defined on stack contents:

$$\bar{z}' u_1 \cdots u_N z_e u_1 \cdots u_N \bar{z} \rightarrow_{\text{skip}} \bar{z}' u_1 \cdots u_N \bar{z}$$

for all  $\bar{z}', u_1, \dots, u_N, z_e, \bar{z} \in \bar{I}^*$  and  $e \in \text{Idem}(\mathbb{M}) \setminus \{\mathbf{0}_{\mathbb{M}}\}$  such that  $\varphi(u_1) = \dots = \varphi(u_N) = \varphi(z_e) = e$  (where the  $u_i$  are in one-to-one correspondence with  $N$ ; we use the subscript  $N$  instead of  $|N|$  for convenience). We then extend it to terms naturally:  $(A, X)[\bar{z}] \rightarrow_{\text{skip}} (A, X)[\bar{z}']$  whenever  $\bar{z} \rightarrow_{\text{skip}} \bar{z}'$ . Observe that the skip rule erases symbols (*potentially deep*) inside the stack, not just at the top. Nevertheless, we will see that allowing the skip rule does not extend the language beyond its downward closure.

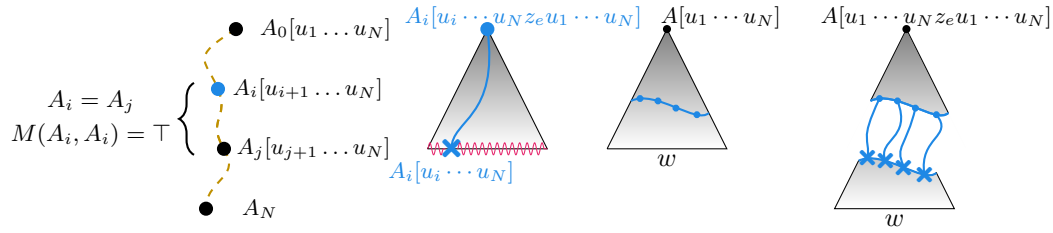
Both the pump rule and skip rules are extended to sentential forms as expected:  $u \rightarrow_{\text{pump}} u'$  if  $u'$  is obtained by applying  $\rightarrow_{\text{pump}}$  to a term in  $u$ , and the same goes for  $\rightarrow_{\text{skip}}$ . We write  $\Rightarrow_{\text{pump, skip}, \bar{\mathcal{G}}}$  for the union of the three relations  $\Rightarrow_{\bar{\mathcal{G}}}$ ,  $\rightarrow_{\text{pump}}$  and  $\rightarrow_{\text{skip}}$ , and  $\stackrel{*}{\Rightarrow}_{\text{pump, skip}, \bar{\mathcal{G}}}$  for its reflexive transitive closure. We may thus define

$$\mathsf{L}_{\text{pump, skip}}(\bar{\mathcal{G}}) = \{w \in T^* \mid (S, U) \stackrel{*}{\Rightarrow}_{\text{pump, skip}} [\bar{\mathcal{G}}]w\}.$$

**Pump and skip are harmless.** We now show that these additional rules do not extend our language beyond its downward closure:

► **Proposition 6.1.**  $\mathsf{L}_{\text{pump, skip}}(\bar{\mathcal{G}}) \subseteq \mathsf{L}(\bar{\mathcal{G}})_{\downarrow}$

The full proof is presented in the full version [47], we sketch the ideas here. For the pump rule this is not hard. It follows from the definition of the product of  $\mathbb{M}$  that if  $e = (B, X, M, A, X)$  is an idempotent, then  $B \mathbf{R}_X A$ , i.e., there must be a derivation from  $(B, X)$  to some  $u(A, X)v$ . By erasing  $u$  and  $v$  (thanks to downward closure and productiveness), from  $(B, X)[\bar{z}]$  we can reach  $(A, X)[\bar{z}]$ , whence we can reach  $(B, X)[z_e \bar{z}]$  (where  $\varphi(z_e) = e$ ) thanks to Lemma 5.2. Hence, any terminal word derived from  $(B, X)[z_e \bar{z}]$  is a subword of a terminal word derived from  $(B, X)[\bar{z}]$  (note that the productiveness of  $\bar{\mathcal{G}}$  is essential here).



■ **Figure 1** The idea behind the skip rule is that if we have a derivation from some non-terminal  $A$  in which we pop  $N$  consecutive infixes  $u_1 \dots u_N$  mapping to the same idempotent  $e = (B, X, M, A, X)$ , then along every branch we must have a node of the form  $A_i[u_{i+1} \dots u_N]$  with  $M(A_i, A_i) = \top$ . This means that for any word mapping to  $e$ , there is a derivation from  $A_i$  popping this word, and with the resulting sentential form containing  $A_i$ . The rest of the sentential form can be erased since we are interested in the downward closure. This non-terminal can be used to “skip” the infix  $u_{i+1} \dots u_N z_e u_1 \dots u_i$ , which maps to  $e$ . Hence, we can turn a derivation from  $A[u_1 \dots u_N]$  into one from  $A[u_1 \dots u_N z_e u_1 \dots u_N]$  by popping the right infix along every branch.

**Eliminating skip.** Replacing applications of the skip rule requires more work. In an indexed grammar, a symbol is pushed once on the stack, but may be popped on multiple branches of a derivation. To tackle this issue, we make the following observations, illustrated by Figure 1. When a sequence  $u \in \bar{I}^*$  with  $\varphi(u) = (B, Y, M, A, X)$  is popped along a branch starting with a non-terminal  $(D, Y)$ , the resulting non-terminal  $(C, X)$  (after popping  $u$ ) must satisfy  $M(D, C) = \top$  (see Lemma 5.3).

Now suppose we are popping a sequence of infixes  $u_1 \dots u_N$ , with  $\varphi(u_1) = \dots = \varphi(u_N) = e = (B, X, M, A, X) \in \text{Idem}(\mathbb{M})$ , along a branch starting with the non-terminal  $(A_0, X)$ . Consider, for  $i = 1, \dots, |N|$ , the non-terminal  $(A_i, X)$  obtained along that branch right after popping  $u_1 \dots u_i$ . By Lemma 5.3, it follows that  $M(A_i, A_{i+1}) = \top$  for  $i = 0, \dots, |N|$ . Since  $e$  is idempotent,  $MM = M$ , and we can then apply the following elementary lemma.

► **Lemma 6.2.** *Let  $M \in \mathbb{B}^{N \times N}$  a matrix. If  $MM = M$  and we have terms  $A_0, \dots, A_N \in N$  such that  $M(A_i, A_{i+1}) = \top$  for all  $i = 1, \dots, |N| - 1$ , then there exists  $i$  with  $M(A_i, A_i) = \top$ .*

**Proof.** By the pigeonhole principle, there exist  $i < j$  such that  $A_i = A_j$ . Since  $M$  is idempotent,  $M^{j-i} = M$ . Hence, we have

$$M(A_i, A_i) = M^{j-i}(A_i, A_j) = \bigwedge_{l=i}^{j-1} M(A_l, A_{l+1}) = \top. \quad \blacktriangleleft$$

Thus, one of the  $A_i$  is such that  $M(A_i, A_i) = \top$  (note that  $i$  may depend on the branch). This means that for all  $\bar{z} \in \bar{I}^*$  with  $\varphi(\bar{z}) = e$ , we have a derivation from  $(A_i, X)[\bar{z}]$  to  $u(A_i, X)u'$  for some  $u, u' \in \text{SF}$  which can be erased since we are only interested in the downward closure (and because of the productiveness property). Hence,  $(A_i, X)$  can “erase” an arbitrary sequence of infixes mapping to  $e$  at the top of its stack. The skip rule is obtained, roughly speaking, by erasing the sequence  $u_{i+1} \dots u_N z_e u_1 \dots u_i$  at the appropriate  $(A_i, X)$  in each branch of a derivation. Since we are sure to encounter, when popping a sequence of  $|N|$  infixes mapping to  $e$ , such a non-terminal on every branch of the derivation (i.e. which can pop such infixes at will), we are able to eliminate all applications of the skip rule.

## 7 Summarizing stack contents

We have shown that adding the pump rule and skip rule to our annotated indexed grammar does not change its downward closure. These rules provide flexibility on stack infixes that map to idempotents: the pump rule can extend them, and the skip rule can reduce them.

We wish to compress stack contents into bounded summaries, where such sequences of infixes are abstracted away, by remembering only the first and last  $|N|$  of them. A natural candidate for this task is Simon's factorization forest theorem [55]. It gives us a way to evaluate a word in a finite monoid via a tree of bounded height, using binary products and arbitrary iterations of idempotent elements. If we only care about the first and last  $|N|$  elements in a sequence of idempotent infixes, we can cut out from this tree all the nodes corresponding to the remaining idempotent infixes.

Two problems remain: First, Simon's theorem gives a linear bound for the height of the tree in the monoid size. In our case this would yield summaries of doubly exponential size, and a quadruply exponential upper bound on an NFA for the downward closure, while our lower bound is only triply exponential. To solve this, we dive deeper into the structure of the monoid, utilizing results by Jecker [37] which let us cut our monoid into polynomially many layers. Within each layer, we decompose the word by reading it from right to left and compressing sequences of idempotent infixes.

The other problem is that to simulate the indexed grammar with a context-free grammar, we use a version of the push operation on the compressed words: from the compressed version of a word  $z$  and a letter  $x$ , we need to be able to compute the compressed version of  $xz$ . This is not a property of Simon's theorem, at least not in its original formulation. One way to circumvent this problem was through the introduction of *forward Ramsey splits* [21], a weaker version of factorizations which still detects sequences of idempotent infixes. In fact, Jecker's results have been applied to improve computational bounds on those splits [45]. It may be interesting to see if one can obtain a form of summaries from such splits, by losing information to obtain a bounded object. Here, however, we do not rely on these, but provide an elementary construction computed deterministically by reading the word right-to-left.

We rely on a recent construction of such summaries presented in [28]. Our exposition is self-contained, however, since we use different notation and our summaries need to be constructed in a slightly different manner. Specifically, theirs need to remember only the first and last infixes, while we need the last  $|N|$ , and theirs are constructed as trees of bounded height, bottom-up, while we need to build ours from right to left along the stack content.

**Green's relations.** We begin with some notions from semigroup theory. Let  $(\mathbf{M}, \cdot, \mathbf{1}_{\mathbf{M}})$  be a finite monoid. We define the usual Green relations  $\mathcal{J}, \mathcal{L}, \mathcal{R}, \mathcal{H}$  on  $\mathbf{M}$ , starting with the following quasi-orders:

- $x \leq_{\mathcal{J}} y$  if there exist  $a, b \in \mathbf{M}$  such that  $x = a \cdot y \cdot b$
- $x \leq_{\mathcal{L}} y$  if there exist  $a \in \mathbf{M}$  such that  $x = a \cdot y$
- $x \leq_{\mathcal{R}} y$  if there exist  $b \in \mathbf{M}$  such that  $x = y \cdot b$
- $x \leq_{\mathcal{H}} y$  if  $x \leq_{\mathcal{L}} y$  and  $x \leq_{\mathcal{R}} y$

The relations  $\mathcal{J}, \mathcal{L}, \mathcal{R}, \mathcal{H}$  are the equivalence relations induced by those quasi-orders: for each  $\mathcal{X} \in \{\mathcal{J}, \mathcal{L}, \mathcal{R}, \mathcal{H}\}$  we define  $\mathcal{X} = \leq_{\mathcal{X}} \cap \geq_{\mathcal{X}}$ , as well as  $<_{\mathcal{X}} = \leq_{\mathcal{X}} \setminus \geq_{\mathcal{X}}$ . We do not include the monoid  $\mathbf{M}$  in the notation since it will always be clear from the context.



**$\mathcal{J}$ -length and  $\mathcal{J}$ -depth.** The *regular  $\mathcal{J}$ -length*<sup>4</sup> of  $\mathbf{M}$ , denoted  $\mathcal{J}\text{len}(\mathbf{M})$ , is defined as

$$\mathcal{J}\text{len}(\mathbf{M}) = \sup\{k \in \mathbb{N} \mid \exists e_1 >_{\mathcal{J}} \cdots >_{\mathcal{J}} e_k \in \text{Idem}(\mathbf{M})\}.$$

► **Definition 7.1.** The *regular  $\mathcal{J}$ -depth* (or just *depth*) of an element  $x \in \mathbf{M}$  is the maximal number  $d$  such that there are  $e_1, \dots, e_d \in \text{Idem}(\mathbf{M})$  with  $e_1 >_{\mathcal{J}} \cdots >_{\mathcal{J}} e_d >_{\mathcal{J}} x$ .

We write  $\text{depth}(x)$  for the regular  $\mathcal{J}$ -depth of an element  $x \in \mathbf{M}$ . We extend this notation to sequences of elements: For  $u \in \mathbf{M}^*$ , we write  $\text{depth}(u)$  instead of  $\text{depth}(\psi_{\mathbf{M}}(u))$ .

► **Remark 7.2.** Note that for all  $x, y \in \mathbf{M}$ , one has  $\text{depth}(x \cdot y) \geq \max(\text{depth}(x), \text{depth}(y))$ . This is because of the definition of  $\mathcal{J}$ , since  $x \cdot y \leq_{\mathcal{J}} x$  and  $x \cdot y \leq_{\mathcal{J}} y$ .

► **Remark 7.3.** Observe that in the case of the stack monoid  $\mathbb{M}$ , we have  $x <_{\mathcal{J}} \mathbf{1}_{\mathbb{M}}$  for every  $x \in \mathbb{M} \setminus \{\mathbf{1}_{\mathbb{M}}\}$ , because if  $y, z \in \mathbb{M} \setminus \{\mathbf{1}_{\mathbb{M}}\}$ , then  $y \cdot z \neq \mathbf{1}_{\mathbb{M}}$ . As a consequence,  $\text{depth}(\mathbf{1}_{\mathbb{M}}) = 0$ , whereas  $\text{depth}(x)$  is in  $[1, \mathcal{J}\text{len}(\mathbb{M})]$  for all  $x \in \mathbb{M} \setminus \{\mathbf{1}_{\mathbb{M}}\}$ . The upper bound follows from the definition of  $\mathcal{J}\text{len}(\mathbb{M})$ . Positivity of  $\text{depth}(x)$  is due to  $x <_{\mathcal{J}} \mathbf{1}_{\mathbb{M}}$ .

**Summaries.** We now define the main object of this section, *summaries*. As mentioned above, they can be viewed as compressions of stack words (with some information loss). Syntactically, these are sequences of (sequences of (...)) sequences, where the nesting depth depends on the depth of  $\varphi(w)$ , where  $w$  is the stack word to be compressed. These sequences contain letters from  $\bar{I}$ , but also a special letter  $e^+$  for each idempotent  $e \in \text{Idem}(\mathbb{M}) \setminus \{\mathbf{1}_{\mathbb{M}}\}$ . Intuitively,  $e^+$  represents a sequence of infixes that evaluate to  $e$ .

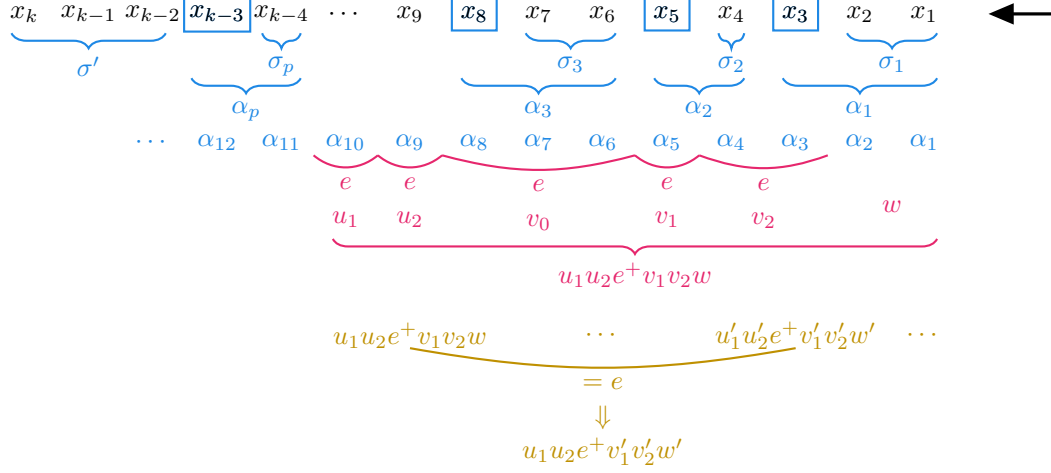
Summaries will have a well-defined image under  $\varphi$ , and the summary of  $w \in \bar{I}^*$  will agree with  $w$  under  $\varphi$ . To this end, we set  $\varphi(e^+) = e$  for all such  $e$ . We also extend  $\varphi$  to  $(\bar{I}^*)^*$  by defining the image of a sequence of words  $z_1 \dots z_n \in (\bar{I}^*)^*$  as the image of their concatenation, and so on for sequences of (sequences of (...)) sequences. To simplify notation, given a sequence of stack symbols  $z \in \bar{I}^*$ , we also write  $\text{depth}(z)$  for  $\text{depth}(\varphi(z))$ , and we extend this notation to sequences of sequences of (sequences of (...)) sequences naturally.

Formally, a 0-summary is just the empty word, and a  $(d+1)$ -summary is a sequence of (i) elements of  $\bar{I} \cup \{e^+ \mid e \in \text{Idem}(\mathbb{M})\}$  and of (ii)  $d$ -summaries. As explained above,  $\varphi(\sigma)$  and  $\text{depth}(\sigma)$  are well-defined for all summaries  $\sigma$ .

► **Definition 7.4.** A *0-summary* is simply the empty word  $\varepsilon$ . Let  $d \in [1, \mathcal{J}\text{len}(\mathbb{M})]$ .

- A  *$d$ -atom* is a word  $(f, X)\sigma$  with  $(f, X) \in \bar{I}$  and  $\sigma$  a  $d'$ -summary for some  $d' < d$ , such that  $\text{depth}((f, X)\sigma) = d$ .
  - A  *$d$ -block* is a sequence of the form  $u_1 \cdots u_N e^+ v_1 \cdots v_N w$  where  $u_1, \dots, u_N, v_1, \dots, v_N$  and  $w$  are sequences of  $d$ -atoms and  $e \in \text{Idem}(\mathbb{M}) \setminus \{\mathbf{1}_{\mathbb{M}}\}$ , such that  $\varphi(u_i) = \varphi(v_i) = e$  for all  $i$  and  $\text{depth}(u_1 \cdots u_N e^+ v_1 \cdots v_N w) = d$ .
  - A  *$d$ -summary* is a sequence of the form  $\sigma' u B_1 \dots B_k$  where  $\sigma'$  is a  $d'$ -summary for some  $d' < d$ ,  $u$  is a word of  $d$ -atoms, and  $B_1, \dots, B_k$  are  $d$ -blocks with  $\text{depth}(\sigma' u B_1 \dots B_k) = d$ .
- A *summary* is a  $d$ -summary for some  $d$ . Observe that by definition, a  $d$ -summary  $\sigma$  has  $\text{depth}(\sigma) = d$ . The set of summaries is denoted **Summaries**.

<sup>4</sup> Called regular  $\mathcal{D}$ -length in [37]. For finite monoids,  $\mathcal{D} = \mathcal{J}$ , and we use  $\mathcal{J}$  here since it is more common. Furthermore, it is defined differently in [37], but a proof that both definitions are equivalent can be found in the long version of that paper [38, Appendix B].



■ **Figure 2** A visual of how summaries are constructed. Say we are given a word  $z$  of regular  $\mathcal{J}$ -depth  $d$ . We read it from right to left. We cut it into  $d$ -atoms by iteratively taking the smallest suffix of regular  $\mathcal{J}$ -depth  $d$ . Its summary consists of its leftmost letter and a  $d'$ -summary for its tail of depth  $d' < d$ . In parallel, we read the resulting sequence of atoms  $\dots \alpha_3 \alpha_2 \alpha_1$  from right to left. Every time we find an infix with a prefix made of  $2|N| + 1$  infixes mapping to the same idempotent, we turn it into a  $d$ -block  $u_1 u_2 e^+ v_1 v_2$ . Finally, whenever we have two blocks corresponding to an idempotent  $e$  and such that the infix between their middle parts also evaluates to  $e$ , we merge them into one  $d$ -block.

**Intuition.** Let us give some intuition on the definition of summaries. Note that when processing a string from right to left, the depth of the growing suffix increases monotonically (Remark 7.2). We read the word from right to left since this is how our stacks are built. Here, a  $d$ -atom represents a suffix-minimal sequence of depth  $d$ : It has depth  $d$ , but when removing the left-most letter, the remainder has lower depth. Hence, that remainder is given as a  $d'$ -summary for some  $d' < d$ . A  $d$ -block can be thought of as a sequence of atoms where some of them (which mapped to  $e$  under  $\varphi$ ) were (lossily) compressed into a single letter  $e^+$ : This compression is only allowed if the compressed part had been surrounded by  $|N|$  sequences of  $d$ -atoms (on each side) that also map to  $e$ , which are still present in  $u_1, \dots, u_N, v_1, \dots, v_N$ . Finally, a  $d$ -summary consists of  $d$ -blocks  $B_1, \dots, B_k$ , some remaining  $d$ -atoms (that did not permit compression into  $d$ -blocks), and a lower-depth part represented by a  $d'$ -summary.

**Compressing stacks into summaries.** Let us describe how a stack  $\bar{z} \in \bar{I}^*$  is compressed into its (unique) summary. For a sequence (i.e. word over a finite alphabet or a summary) of length  $\geq 1$ , its *tail* is obtained by removing its leftmost element. Given a word  $\bar{z}$  with  $\text{depth}(\bar{z}) = d$ , we proceed in three stages, illustrated in Figure 2.

*Stage I: Splitting into  $d$ -atoms.* First we split the word into suffix-minimal infixes of depth  $d$ , from right to left, and a residual prefix of some depth  $d' < d$ . By suffix-minimality, the tail of each of these depth- $d$  infixes has depth  $< d$ . We thus turn each of these depth- $d$  infixes into a  $d$ -atom by computing recursively a summary for its lower-depth tail. Afterwards, the residual prefix of depth  $d' < d$  is recursively turned into a summary.

*Stage II: Compression into blocks.* Then, we look at the sequence of  $d$ -atoms and their values in  $\mathbb{M}$ . Going from right to left, we look for sequences of  $2|N| + 1$  infixes  $u_1 \dots u_N v_0 v_1 \dots v_N$  all evaluating to the same idempotent. Whenever we find such a pattern we turn the current infix (i.e. the pattern with, possibly, a suffix  $w$ ) into a  $d$ -block by replacing the middle part  $v_0$  by  $e^+$ . There may be a remainder  $u$  at the left end of the  $d$ -atom sequence with no such pattern; we write it explicitly in the summary.

*Stage III: Merging blocks.* Finally, we look at the sequence of resulting  $d$ -blocks. We go again from right to left, this time looking for pairs of blocks corresponding to the same idempotent  $e$ , and so that the infix between their  $e^+$  markers also evaluates to  $e$ . This lets us merge them into a single block, obtained by abstracting their  $e^+$  markers and everything in between in a single  $e^+$ .

**Updating a summary.** A key feature of summaries is that given the summary  $\sigma$  of a stack  $\bar{z} \in \bar{I}^*$  and an additional letter  $\bar{x} \in \bar{I}$ , we can compute the summary of  $\bar{x}\bar{z}$ . This is needed when simulating pushes. We now describe the relevant operation,

$$\text{push}(\_ \blacktriangleright \_) : \bar{I}^* \times \text{Summaries} \rightarrow \text{Summaries}.$$

Given a depth  $d \in [0, \mathcal{J}\text{len}(\mathbb{M})]$ , an element  $(f, X) \in \bar{I}$  and an  $d$ -summary  $\sigma$ , we define  $\text{push}((f, X) \blacktriangleright \sigma)$  as follows. If  $d = 0$  then  $\sigma = \varepsilon$  and  $\text{push}((f, X) \blacktriangleright \sigma) = (f, X)$ . If  $d > 0$  then  $\sigma$  is of the form  $\sigma' u B_1 \dots B_k$ .

1. If  $\text{depth}((f, X)\sigma) > d$  then let  $d_+ = \text{depth}((f, X)\sigma)$ . We set  $\text{push}((f, X) \blacktriangleright \sigma)$  as the  $d_+$ -summary made of a single  $d_+$ -atom  $(f, X)\sigma$
2. Otherwise  $\text{depth}((f, X)\sigma) \leq d$ . In fact, since  $\text{depth}((f, X)\sigma)$  is at least  $\text{depth}(\sigma) = d$ , we even know  $\text{depth}((f, X)\sigma) = d$ .
  - a. If  $\text{depth}((f, X)\sigma') < d$  then
$$\text{push}((f, X) \blacktriangleright \sigma) = (\text{push}((f, X) \blacktriangleright \sigma')) u B_1 \dots B_k$$
  - b. Otherwise,  $\text{depth}((f, X)\sigma') \geq d$ . Then we even know  $\text{depth}((f, X)\sigma') = d$ , because  $\text{depth}((f, X)\sigma')$  is at most  $\text{depth}((f, X)\sigma) \leq d$ . Hence,  $(f, X)\sigma'$  is a  $d$ -atom.
    - i. If the sequence of  $d$ -atoms  $((f, X)\sigma')u$  is of the form
$$u_1 \dots u_N v_0 v_1 \dots v_N w \text{ with } \varphi(v_0) = \varphi(u_i) = \varphi(v_i) = e \text{ for all } i \geq 1, \text{ for some } e \in \text{Idem}(\mathbb{M}), \text{ then define the } d\text{-block } B = u_1 \dots u_N e^+ v_1 \dots v_N w \text{ (note that } e^+ \text{ replaces the middle infix } v_0).$$
      - A. Suppose there is  $j$  such that  $B_j$  is of the form

$$u'_1 \dots u'_N e^+ v'_1 \dots v'_N w'$$

and the infix  $v_1 \dots v_N w B_1 \dots B_{j-1} u'_1 \dots u'_N$  also evaluates to  $e$  under  $\varphi$ . In this case, we merge everything up to  $B_j$  into a single block: We define  $\text{push}((f, X) \blacktriangleright \sigma)$  to be  $B' B_j B_{j+1} \dots B_k$ , with  $B' = u_1 \dots u_N e^+ v'_1 \dots v'_N w'$ . When there are multiple such  $j$ , we pick the maximal one.

B. Otherwise  $\text{push}((f, X) \blacktriangleright \sigma) = B B_1 \dots B_k$

- ii. Otherwise  $\text{push}((f, X) \blacktriangleright \sigma) = u' B_1 \dots B_k$  with  $u' = ((f, X)\sigma')u$

The definition is extended inductively to  $\bar{I}^*$  by

$$\text{push}((f, X)\bar{z} \blacktriangleright \sigma) = \text{push}((f, X) \blacktriangleright \text{push}(\bar{z} \blacktriangleright \sigma)).$$

Intuitively, the  $\text{push}(\_ \blacktriangleright \_)$  operation simultaneously executes stages I, II and III described above. Upon adding a new letter  $(f, X)$ , it checks whether it yields a new atom (Stage I). Case (1) happens when a new atom is generated because the depth of the summary is increased by adding  $(f, X)$ , case (a) when no new atom is generated at depth  $d$ . In case (b), we have a new atom  $(f, X)\sigma'$ . We thus have to check whether this new atom yields a new block (Stage II). If it does, we are in case (i), and we then have to apply Stage III, i.e., see if we can merge this new block with another one. In case (A) we can; in case (B) we cannot.

**Popping from summaries.** We also define the inverse relation: for  $\bar{z} \in \bar{I}^*$  and  $\sigma, \sigma' \in \text{Summaries}$ , we have  $\sigma \in \text{pop}(\bar{z} \blacktriangleleft \sigma')$  if and only if  $\sigma' = \text{push}(\bar{z} \blacktriangleright \sigma)$ . Note that  $\text{push}(\_ \blacktriangleright \_)$  is a function while  $\text{pop}(\_ \blacktriangleleft \_)$  is a relation. Intuitively,  $\text{push}(\_ \blacktriangleright \_)$  makes compressions, by creating and merging blocks, thereby losing information. Meanwhile,  $\text{pop}(\_ \blacktriangleleft \_)$  may revert those compressions, and thus requires non-determinism.

**Bounding summaries.** The *size* of a  $d$ -summary (resp.  $d$ -atom,  $d$ -block) is defined recursively as the sum of the *sizes* of its elements, the size of a single letter being 1. Although summaries can have a priori unbounded sizes, the ones we will use in our context-free grammar will be of the form  $\text{push}(\bar{z} \blacktriangleright \varepsilon)$  for some  $\bar{z} \in \bar{I}^*$ . For such summaries, we can prove a size bound:

► **Lemma 7.5.** *For each  $\bar{z} \in \bar{I}^*$ , the size of  $\text{push}(\bar{z} \blacktriangleright \varepsilon)$  is at most exponential in  $|N|$ .*

To prove this, we rely on two results of Jecker [37], recalled below.

Define the *Ramsey function* of  $\mathbf{M}$  as follows: for all  $k \in \mathbb{N}$ ,  $\mathcal{R}_{\mathbf{M}}(k)$  is the minimal  $n$  such that for every word of length  $n$  there exists  $e \in \text{Idem}(\mathbf{M})$  and  $u_1, \dots, u_k \in \mathbf{M}^+$  such that the word contains an infix  $u_1 \cdots u_k$  with  $\psi_{\mathbf{M}}(u_i) = e$  for all  $i$ .

The first one guarantees that in a sequence of exponential length in  $\mathcal{J}\text{len}(\mathbf{M})$ , we can find large sequences of consecutive infixes that all map to the same idempotent [37, Theorem 1]:

► **Theorem 7.6** (Jecker). *For all finite monoid  $\mathbf{M}$ , for all  $k \in \mathbb{N}$ :*

$$\mathcal{R}_{\mathbf{M}}(k) \leq (k|\mathbf{M}|^4)^{\mathcal{J}\text{len}(\mathbf{M})}.$$

The other one bounds the regular  $\mathcal{J}$ -length of a Boolean matrix monoid by a polynomial in its dimension [37, Theorem 2]:

► **Theorem 7.7** (Jecker). *The regular  $\mathcal{J}$ -length of  $\mathbb{B}^{N \times N}$  is  $\mathcal{J}\text{len}(\mathbb{B}^{N \times N}) = \frac{N^2 + N + 2}{2}$ .*

For Lemma 7.5, we show that when viewing  $d$ -atoms and  $d'$ -summaries, for  $d' < d$ , as individual letters, the size of a  $d$ -summary is bounded by an exponential in  $|N|$ . The overall bound on summaries then results from the product of those (polynomially many) exponential functions.

## 8 Building the context-free grammar

We now construct a context-free grammar  $\mathcal{C}_{\mathcal{G}}$  that over-approximates  $\mathbf{L}(\mathcal{G})$ , but remains within its downward closure. It will thus satisfy  $\mathbf{L}(\mathcal{C}_{\mathcal{G}}) \downarrow = \mathbf{L}(\mathcal{G}) \downarrow$  and enable us to compute an NFA for the latter.

**Definition of the grammar.** Essentially,  $\mathcal{C}_{\mathcal{G}}$  is obtained from  $\bar{\mathcal{G}}$  by replacing stack contents with their summaries.

We first restrict the set of summaries to those that can actually result from a derivation of the indexed grammar. A summary  $\sigma$  is *feasible* if  $\sigma = \text{push}(\bar{z} \blacktriangleright \varepsilon)$  for some feasible  $\bar{z} \in \bar{I}^*$ . Note that this implies  $\varphi(\sigma) \neq \mathbf{0}_{\mathbf{M}}$  by Lemma 5.1. The non-terminals will be triples  $(A, X, \sigma)$  where  $(A, X) \in \bar{N}$ , and  $\sigma$  is a feasible summary. We call such triples *feasible*. The set of feasible triples is denoted **FT**.

Define the following context-free grammar  $\mathcal{C}_{\mathcal{G}}$ : Its set of non-terminals is **FT**, with  $(S, \mathbf{U}, \varepsilon)$  the initial one. The set of terminal symbols is  $T$ . The productions directly mimic the productions in  $\bar{\mathcal{G}}$ , except that push and pop productions are simulated by the  $\text{push}(\_ \blacktriangleright \_)$  and  $\text{pop}(\_ \blacktriangleleft \_)$  relations on summaries:

- If  $(A, X) \rightarrow w \in \bar{P}$  then  $(A, X, \sigma) \rightarrow w$ , for all feasible  $\sigma$ .
- If  $(A, X) \rightarrow (B, X)(C, X) \in \bar{P}$  then  $(A, X, \sigma) \rightarrow (B, X, \sigma)(C, X, \sigma)$ , for all feasible  $\sigma$ .
- If  $(A, X) \rightarrow (B, Y)(f, X) \in \bar{P}$  then  $(A, X, \sigma) \rightarrow (B, Y, \text{push}((f, X) \blacktriangleright \sigma))$ , for all summary  $\sigma$  so that  $\sigma$  and  $\text{push}((f, X) \blacktriangleright \sigma)$  are feasible.
- If  $(A, Y)(f, X) \rightarrow (B, X) \in \bar{P}$  then  $(A, Y, \sigma) \rightarrow (B, X, \sigma')$ , whenever  $\sigma, \sigma'$  are feasible and  $\sigma' \in \text{pop}((f, X) \blacktriangleleft \sigma)$ .

Abusing terminology slightly, we call production rules of the third and fourth type *pushes* and *pops*, respectively (even though they are standard context-free productions).

**Correctness of the construction.** The key property of  $\mathcal{C}_{\mathcal{G}}$  is that it has the same downward closure as  $\mathcal{L}(\mathcal{G})$ .

► **Theorem 8.1.**  $\mathcal{L}(\mathcal{C}_{\mathcal{G}})\downarrow = \mathcal{L}(\mathcal{G})\downarrow$ .

One of the directions is quite easy: we can simply show that the language of  $\mathcal{C}_{\mathcal{G}}$  contains that of  $\bar{\mathcal{G}}$ . This is natural as  $\mathcal{C}_{\mathcal{G}}$  is built as an over-approximation of  $\bar{\mathcal{G}}$ . We can turn a derivation of  $\bar{\mathcal{G}}$  into one of  $\mathcal{C}_{\mathcal{G}}$  by replacing every stack content with its summary.

► **Proposition 8.2.**  $\mathcal{L}(\bar{\mathcal{G}}) \subseteq \mathcal{L}(\mathcal{C}_{\mathcal{G}})$ .

**Simulating  $\mathcal{C}_{\mathcal{G}}$  with pumps and skips.** For the inclusion  $\mathcal{L}(\mathcal{C}_{\mathcal{G}})\downarrow \subseteq \mathcal{L}(\mathcal{G})\downarrow$ , we will show that every derivation in  $\mathcal{C}_{\mathcal{G}}$  can be simulated by pumps and skips:

► **Proposition 8.3.**  $\mathcal{L}(\mathcal{C}_{\mathcal{G}}) \subseteq \mathcal{L}_{\text{pump,skip}}(\bar{\mathcal{G}})$ .

Indeed, by Proposition 6.1, this implies that  $\mathcal{L}(\mathcal{C}_{\mathcal{G}}) \subseteq \mathcal{L}(\bar{\mathcal{G}})\downarrow$  and thus  $\mathcal{L}(\mathcal{C}_{\mathcal{G}})\downarrow \subseteq \mathcal{L}(\bar{\mathcal{G}})\downarrow = \mathcal{L}(\mathcal{G})\downarrow$ , establishing Theorem 8.1.

We shall prove Proposition 8.3 by simulating summaries by actual stacks. Recall that in the summary  $\text{push}(\bar{z} \blacktriangleright \varepsilon)$  that compresses the stack  $\bar{z}$ , the letters  $e^+$  represent a sequence of  $e$ -words, where we call  $\bar{w} \in \bar{I}^*$  an *e-word* if  $\varphi(\bar{w}) = e$ . The other letters in  $\sigma$  are taken directly from  $\bar{z}$ . Therefore, the unfolding reverses this: It replaces  $e^+$  by  $e$ -words and leaves the other letters unchanged.

Intuitively, our strategy for replacing  $e^+$  is as follows. We replace  $e^+$  by a sequence of all  $e$ -words that might be needed to sustain the remainder of the derivation (specifically, all its pop operations). In a particular branch of the derivation, those  $e$ -words that are not needed can always be cleared using the skip rule. On the other hand, the pump rule allows us to justify introducing all these  $e$ -words.

**Unfoldings.** Instead of tailoring the stack  $\bar{z}$  simulating  $\sigma$  to a specific derivation, we will construct a “canonical” stack word  $\bar{z}$  for  $\sigma$  that only depends on the number of pops in that derivation. We will call this canonical stack word the “ $p$ -unfolding” of  $\sigma$ , where  $p$  is the number of pops it is designed to sustain. This means, intuitively, each  $e^+$  is replaced by a large enough concatenation of  $e$ -words so that any sequence of  $p$  pops can be executed.

Just to choose the order in which those  $e$ -words appear, we impose an arbitrary total order (e.g. a length-lexicographical ordering) on the set of summaries, which we denote  $\triangleleft$ .

► **Definition 8.4.** Let  $p \in \mathbb{N}$  and  $\sigma$  be a  $d$ -summary. The  *$p$ -unfolding* of  $\sigma$ , denoted  $\text{unf}_p(\sigma)$  is defined inductively w.r.t.  $d$  and  $p$  as follows.

For  $d = 0$ , the  $p$ -unfolding of a 0-summary (i.e.,  $\varepsilon$ ) is  $\varepsilon$ .

- The  $p$ -unfolding of a  $d$ -atom  $(f, X)\sigma'$  is  $(f, X)\text{unf}_p(\sigma')$ .
- The  $p$ -unfolding of a sequence of  $d$ -atoms  $\alpha_1 \cdots \alpha_m$  is defined as  $\text{unf}_p(\alpha_1) \cdots \text{unf}_p(\alpha_m)$ .
- The  $p$ -unfolding of a  $d$ -block  $B = u_1 \dots u_N e^+ v_1 \dots v_N w$  is

$$z^u z^v ((f, X)\text{unf}_{p-1}(\sigma_1)) \cdots ((f, X)\text{unf}_{p-1}(\sigma_r)) z^v \text{unf}_p(w),$$

where:

- $(f, X)$  is the first symbol of  $B$ , that is, the stack symbol such that the first  $d$ -atom of  $u_1$  is of the form  $(f, X)\sigma'$ .
- $z^u = \text{unf}_p(u_1 \dots u_N)$  and  $z^v = \text{unf}_p(v_1 \dots v_N)$
- if  $p > 0$ , then  $(\sigma_i)_{1 \leq i \leq r}$  is the family of all feasible summaries  $\sigma'$  for which we have the equality  $\text{push}((f, X) \blacktriangleright \sigma') = u_1 \dots u_N e^+ v_1 \dots v_N$ , ordered according to  $\leq$ .
- if  $p = 0$ , then  $r = 0$ , i.e., the  $p$ -unfolding of  $B$  is simply  $z^u z^v \text{unf}_p(w)$ .
- The  $p$ -unfolding of a  $d$ -summary  $\sigma' u B_1 \dots B_m$  is defined as

$$\text{unf}_p(\sigma') \text{unf}_p(u) \text{unf}_p(B_1) \cdots \text{unf}_p(B_m).$$

Since the unfolding is obtained by replacing each  $e^+$  in a summary by a concatenation of words with image  $e$  (and all other letters are unchanged), the unfolding has the same image under  $\varphi$  as  $\sigma$ :

► **Remark 8.5.** For every summary  $\sigma$  and every  $p \in \mathbb{N}$ , we have  $\varphi(\text{unf}_p(\sigma)) = \varphi(\sigma)$ .

**Removing excess  $e$ -words.** When choosing a stack word to simulate a given summary  $\sigma$ , we pick the  $p$ -unfolding, where  $p$  is the total number of pops in the entire derivation. However, some branches will apply less than  $p$  pops. The following lemma is therefore crucial: It allows us to get rid of excess  $e$ -words that are not needed on less pop-heavy branches, while maintaining the invariant that we have unfoldings on the stack:

► **Lemma 8.6.** For each  $\sigma$  and  $p \geq 1$ :  $\text{unf}_p(\sigma) \xrightarrow{*}_{\text{skip}} \text{unf}_{p-1}(\sigma)$ .

Note that it is not possible to simply skip  $e$ -words at will, since that requires equality of some infixes in the stack. However, unfoldings are carefully constructed to allow Lemma 8.6. For example, the fact that we always follow a uniform order  $\leq$  on summaries is key.

**Simulating pushes using unfoldings.** Let us now show how unfoldings are used to mimic derivations of  $\mathcal{C}_{\mathcal{G}}$  in  $\bar{\mathcal{G}}$ , with pumps and skips. First, note that all productions of  $\mathcal{C}_{\mathcal{G}}$  that are not pushes and pops have direct counterparts in  $\bar{\mathcal{G}}$ . Suppose we want to simulate a push, say a rule  $(A, X, \sigma_1) \rightarrow (B, Y, \sigma_2)$  of  $\mathcal{C}_{\mathcal{G}}$ , induced by a push rule  $(A, X) \rightarrow (B, Y)(f, X)$  in  $\bar{\mathcal{G}}$ . If  $(A, X, \sigma_1)$  is simulated by  $(A, X)[\text{unf}_p(\sigma_1)]$ , then we can use  $(A, X) \rightarrow (B, Y)(f, X)$  first in  $\bar{\mathcal{G}}$ . But then the stack is  $(f, X)\text{unf}_p(\sigma_1)$ , rather than an unfolding of  $\sigma_2$ . The following lemma tells us that, using pump and skip, we can replace  $(f, X)\text{unf}_p(\sigma_1)$  by  $\text{unf}_p(\sigma_2)$ , which will then enable us to continue the simulation.

► **Lemma 8.7.** Let  $p \in \mathbb{N}$ , let  $(A, X, \sigma_1) \rightarrow (B, Y, \sigma_2)$  a rule of  $\mathcal{C}_{\mathcal{G}}$  with  $\sigma_2 = \text{push}((f, X) \blacktriangleright \sigma_1)$ . We have

$$(B, Y)[(f, X)\text{unf}_p(\sigma_1)] \xRightarrow{*}_{\text{pump, skip}} (B, Y)[\text{unf}_p(\sigma_2)].$$

In fact, only  $\rightarrow_{\text{pump}}$  and  $\rightarrow_{\text{skip}}$  are needed to prove this lemma.



**Simulating pops using unfoldings.** Pop steps are more difficult. In a production  $(A, X, \sigma_1) \rightarrow (B, Y, \sigma_2)$  in  $\mathcal{C}_G$  induced by a pop rule  $(A, X)(f, Y) \rightarrow (B, Y)$ ,  $\sigma_2$  is the summary of a word obtained by removing the first letter of a word compressed by  $\sigma_1$ . This removal might break a block centered around some  $e \in \text{Idem}(\mathbb{M})$ , in  $\sigma_1$ . This means, the symbol  $e^+$  is replaced by a concatenation of summaries. However,  $p$ -unfoldings are designed so that  $\text{unf}_p(\sigma_1)$  contains enough  $e$ -words for each  $e^+$  so that using skip rules, we can remove a subset of them so that the resulting stack is precisely  $(f, Y)\text{unf}_{p-1}(\sigma_2)$ . This is shown in the following lemma:

► **Lemma 8.8.** *Let  $p \geq 1$ , let  $(A, X, \sigma_1) \rightarrow (B, Y, \sigma_2)$  a rule of  $\mathcal{C}_G$  with  $\sigma_2 \in \text{pop}((f, Y) \blacktriangleleft \sigma_1)$ . We have*

$$(A, X)[\text{unf}_p(\sigma_1)] \xrightarrow{*}_{\text{skip}} (A, X)[(f, Y)\text{unf}_{p-1}(\sigma_2)].$$

Thus, to simulate this pop rule, we can first invoke Lemma 8.8 and then apply the production  $(A, X)(f, Y) \rightarrow (B, Y)$ .

**Simulating the whole derivation.** With Lemmas 8.7 and 8.8 in hand, Proposition 8.3 is now easy to show. Indeed, it is straightforward to simulate an entire derivation of  $\mathcal{C}_G$  in  $\bar{\mathcal{G}}$  with pump and skip rules: Simulating pushes and pops is as explained in Lemmas 8.7 and 8.8, and the other productions are immediate.

We have thus completed Proposition 8.3 and hence Theorem 8.1.

**Putting it all together.** With Theorem 8.1, we are prepared to prove Theorem 3.1. By Lemma 7.5, the size of a summary is bounded by an exponential in the size of  $\mathcal{G}$ . As a consequence, the size of  $\mathcal{C}_G$  is at most doubly exponential in the size of  $\mathcal{G}$ . Since for a given context-free grammar, one can compute an exponential-sized NFA for its language's downward closure [12, Corollary 6], this yields a triply exponentially sized NFA for  $\text{L}(\mathcal{G})\downarrow$ , as desired in Theorem 3.1.

## 9 Lower bounds

In this section, we prove the lower bounds in our main results.

**NFA Lower bound: Overview.** We begin with Theorem 3.2. The overall goal is to have a unique complete derivation tree that is a full binary tree of doubly exponential depth: clearly, such a tree must have triply exponentially many leaves. Here, the challenge is to ensure that the paths have doubly exponential length.

A standard construction can enforce *singly* exponentially long paths: Use a height- $n$  stack over the alphabet  $\{0, 1\}$ , and then count up from  $0^n$  to  $1^n$ , resulting in  $2^n - 1$  steps. This works because when incrementing a binary expansion of the form  $1^m 0 w$  (the least significant digit being on the left), we must replace the prefix  $1^m 0$  with  $0^m 1$ . Here, we store the number  $m \leq n$  in the non-terminal, so as to restore the stack height  $n$  when pushing  $0^m 1$ .

Doing the same for stack height  $2^n$  is not so easy: To restore a stack height of  $2^n$ , we would need to remember (in the non-terminal) a number  $m \leq 2^n$ . In fact, enforcing a single run of length  $2^{2^n}$  in a pushdown automaton of polynomial size is not possible: A pushdown automaton that accepts any word must also accept a word of at most exponential length.

Instead, we exploit the fact that an indexed grammar can simulate a pushdown automaton *with alternation*: We implement binary counting on a stack of height  $2^n$ ; in order to replace a prefix  $1^m 0$  with  $0^m 1$ , we non-deterministically push some number of 0's, but then use alternation to ensure that the stack height is exactly  $2^n$ .



**Step I: Checking the stack via alternation.** To this end, we introduce syntactic sugar. We will use rules of the form

$$A \xrightarrow{\mathcal{A}_1, \dots, \mathcal{A}_r} B, \quad (1)$$

where  $\mathcal{A}_1, \dots, \mathcal{A}_r$  are DFAs over the stack alphabet  $I$ . The rule has the same effect as  $A \rightarrow B$ , but it can only be applied to a term  $A[z]$  if for each  $i = 1, \dots, r$ , the stack  $z$  has a prefix in  $L(\mathcal{A}_i)$ . Such rules can be implemented with only polynomial overhead: Introduce non-terminals  $C_i, D_i$  for each  $i = 1, \dots, r$  and also  $E_q$  for each state  $q$  in the DFAs  $\mathcal{A}_1, \dots, \mathcal{A}_r$  (we assume the state sets are disjoint). Then we can simulate (1) using  $A \rightarrow D_1 C_1$ ,  $D_i \rightarrow D_{i+1} C_{i+1}$  for  $i = 1, \dots, r-1$ , and  $D_r \rightarrow B$ , which split the term  $A[z]$  into terms  $B[z]$  and  $C_1[z], \dots, C_r[z]$ . We then run  $\mathcal{A}_i$  using rules  $C_i \rightarrow E_{q_i}$ , where  $q_i$  is the initial state of  $\mathcal{A}_i$ , for each  $i$ . The non-terminals  $E_q$  simulate the DFAs: for each transition  $(p, f, q)$ , we have  $E_p f \rightarrow E_q$ . To check acceptance, we have  $E_q \rightarrow \varepsilon$  for each final state  $q$ .

**Step II: Implementing a binary counter in DFAs.** We want to use rules (1) to check that the current stack height is  $2^n$ , for which we construct automata  $(\mathcal{A}_i)_{1 \leq i \leq n}$  over some alphabet  $\Sigma_n$  such that: (i) Each  $\mathcal{A}_i$  has two states  $\mathbf{0}_i$  and  $\mathbf{1}_i$ , with  $\mathbf{0}_i$  being initial and  $\mathbf{1}_i$  the only final state, (ii)  $|\Sigma_n| = n$ , and (iii) the intersection  $\bigcap_{i=1}^n L(\mathcal{A}_i)$  contains a single word of length  $2^n$ . The construction is simpler if we do this for  $2^n - 1$  instead of  $2^n$ , which suffices: We can introduce a fresh letter  $\#$  and build automata  $\mathcal{A}'_i$  with  $L(\mathcal{A}'_i) = L(\mathcal{A}_i)\#$ .

The idea is simply to use  $\Sigma_n = \{\text{inc}_1, \dots, \text{inc}_n\}$ . Each letter is an increment operation over an  $n$ -bit binary counter:  $\text{inc}_i$  should be read as “flip the  $i$ -th bit from 0 to 1 and all lower bits from 1 to 0”. Each automaton  $\mathcal{A}_i$  keeps track of the value of the  $i$ -th bit throughout that sequence of instructions.

Formally,  $\mathcal{A}_i = (\{\mathbf{0}_i, \mathbf{1}_i\}, \Sigma_n, \delta_i, \mathbf{0}_i, \{\mathbf{1}_i\})$  where  $\delta_i(\mathbf{0}_i, \text{inc}_j)$  is defined as  $\mathbf{1}_i$  if  $j = i$ , it is  $\mathbf{0}_i$  if  $j < i$ , and it is undefined if  $j > i$ . Meanwhile,  $\delta_i(\mathbf{1}_i, \text{inc}_j)$  is  $\mathbf{0}_i$  if  $j > i$ , it is  $\mathbf{1}_i$  if  $j < i$  and it is undefined if  $i = j$ . It is easy to check that there is a unique word accepted by those automata, corresponding to the only correct sequence of instructions to increment a  $n$ -bit binary counter from 0 to  $2^n - 1$ , which enforces a single string of length  $2^n - 1$ , as desired.

**Step III: Constructing the indexed grammar.** Let  $\mathcal{A}_1, \dots, \mathcal{A}_n$  the DFAs over  $\Sigma_n$  built above. Since they will check for exponential stack height, but we also need to store the binary digits on the stack, we modify them slightly. For each  $i$ , the automaton  $\mathcal{B}_i$  will work over the alphabet  $\{\perp\} \cup (\Sigma_n \times \{0, 1\})$  and accept exactly the words of the form  $(\alpha_1, b_1) \cdots (\alpha_m, b_m) \perp$  where  $\alpha_1 \cdots \alpha_m \in L(\mathcal{A}_i)$ . The  $\perp$  letter is used to mark the bottom of the stack.

Consider the following grammar:  $\mathcal{G}_n = (N_n, T, I_n, P_n, S)$  with  $N_n = \{S, A, B, D, F, Z\}$ ,  $T = \{\mathbf{a}\}$ ,  $I_n = \{\perp\} \cup (\Sigma_n \times \{0, 1\})$ , and  $P_n$  contains the following rules:

$$\begin{array}{lll} S \rightarrow Z \perp & D \rightarrow AA & B \rightarrow Z(\alpha, 1) \\ Z \rightarrow Z(\alpha, 0) & A(\alpha, 1) \rightarrow A & A \perp \rightarrow F \\ Z \xrightarrow{\mathcal{B}_1, \dots, \mathcal{B}_n} D & A(\alpha, 0) \rightarrow B & F \rightarrow \mathbf{a} \end{array}$$

for each  $\alpha \in \Sigma_n$ . The grammar works as follows. Initially, it places  $\perp$  on the stack and switches to  $Z$ . A non-terminal  $Z$  will then fill the stack with 0's, each of which is non-deterministically annotated with some  $\alpha \in \Sigma_n$ . After pushing these, it verifies that the stack height is  $2^n$ , by using  $Z \xrightarrow{\mathcal{B}_1, \dots, \mathcal{B}_n} D$ . This  $D$  splits into two  $A$ 's, where an increment is performed on the number encoded on the stack: It removes the prefix of the form  $1^m 0$  and switches to  $B$ . After this, it has to put back  $0^m 1$ : To this end, it pushes a single 1

and then using  $Z$  pushes 0's non-deterministically. It then uses  $Z \xrightarrow{B_1, \dots, B_n} D$  to verify that the stack height is  $2^n$ . All this repeats until in each branch, all stack contents encode the number  $2^{2^n} - 1$ . This means, all terms are of the form  $A[z\perp]$ , where  $z$  has length  $2^n$  and all its digits are 1's. Each such  $A[z\perp]$  is then rewritten to  $F$ , and then to  $\mathbf{a}$ . Since the terms are duplicated before each increment (using  $D \rightarrow AA$ ), the final number of  $\mathbf{a}$  letters is  $\exp_3(n)$ , deriving  $\mathbf{a}^{\exp_3(n)}$ . It is also straightforward to check that  $\mathbf{a}^{\exp_3(n)}$  is the only derivable word.

**Computational hardness.** The lower bounds for downward closure inclusion and equivalence now follow easily from Theorem 3.2 and results in [60]. In [60], the  $\Delta(f)$  property of language classes is introduced. Roughly speaking, it requires simple closure properties and that for given  $n \in \mathbb{N}$ , one can construct in polynomial time the language  $\{\mathbf{a}^{f(n)}\}$ . Under additional mild assumptions, [60, Theorem 15] shows that downward closure inclusion and equivalence are  $\text{coTIME}(f)$ -hard for  $\Delta(f)$  classes. Since all assumptions besides a small grammar for  $\{\mathbf{a}^{\exp_3(n)}\}$  are easy to observe, we may conclude that the indexed languages are  $\Delta(\exp_3)$  and the two problems are  $\text{co-3-NEXP-hard}$ .

**DFA size.** For Theorem 3.5, we adapt an idea from [12, Theorem 7], which shows a doubly exponential lower bound for downward closure DFAs for CFL. It is not difficult to translate the grammar  $\mathcal{G}_n$  for  $\{\mathbf{a}^{\exp_3(n)}\}$  into one for  $L_n = \{uv \mid u, v \in \{0, 1\}^* \mid |u| = |v| = \exp_3(n), u \neq v\}$ . It is easy to see that a DFA for  $L_n \downarrow$  requires  $\exp_4(n)$  states: For distinct  $u, v \in \{0, 1\}^*$  with  $|u| = |v|$ , the DFA must accept  $uv$  and  $vu$ , but reject  $uu$  and  $vv$ . It therefore must enter distinct states after reading  $u$  and  $v$ .

## 10 Conclusion

We have established (asymptotically) tight bounds on the size of an automaton for the downward closure of an indexed language. We rely on an algebraic abstraction of stack contents to translate indexed grammars into context-free ones while preserving the downward closure.

---

## References

- 1 Parosh Aziz Abdulla, Luc Boasson, and Ahmed Bouajjani. Effective lossy queue languages. In Fernando Orejas, Paul G. Spirakis, and Jan van Leeuwen, editors, *Automata, Languages and Programming, 28th International Colloquium, ICALP 2001, Crete, Greece, July 8-12, 2001, Proceedings*, volume 2076 of *Lecture Notes in Computer Science*, pages 639–651. Springer, 2001. doi:10.1007/3-540-48224-5\_53.
- 2 Alfred V. Aho. Indexed grammars - an extension of context-free grammars. *J. ACM*, 15(4):647–671, 1968. doi:10.1145/321479.321488.
- 3 C. Aiswarya, Pascal Baumann, Prakash Saivasan, Lia Schütze, and Georg Zetsche. Bounded treewidth, multiple context-free grammars, and downward closures. *Proc. ACM Program. Lang.*, 10(POPL), January 2026. doi:10.1145/3776716.
- 4 C. Aiswarya, Soumodev Mal, and Prakash Saivasan. On the satisfiability of context-free string constraints with subword-ordering. In Christel Baier and Dana Fisman, editors, *LICS '22: 37th Annual ACM/IEEE Symposium on Logic in Computer Science, Haifa, Israel, August 2 - 5, 2022*, pages 6:1–6:13. ACM, 2022. doi:10.1145/3531130.3533329.
- 5 Antoine Amarilli, Florin Manea, Tina Ringleb, and Markus L. Schmid. Linear time subsequence and supersequence regex matching. In Pawel Gawrychowski, Filip Mazowiecki, and Michal Skrzypczak, editors, *50th International Symposium on Mathematical Foundations of Computer Science, MFCS 2025, Warsaw, Poland, August 25-29, 2025*, LIPIcs, pages 9:1–9:19. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2025. doi:10.4230/LIPIcs.MFCS.2025.9.

- 6 Ashwani Anand, Sylvain Schmitz, Lia Schütze, and Georg Zetsche. Verifying unboundedness via amalgamation. In Pawel Sobocinski, Ugo Dal Lago, and Javier Esparza, editors, *Proceedings of the 39th Annual ACM/IEEE Symposium on Logic in Computer Science, LICS 2024, Tallinn, Estonia, July 8-11, 2024*, pages 4:1–4:15. ACM, 2024. doi:10.1145/3661814.3662133.
- 7 Ashwani Anand and Georg Zetsche. Priority downward closures. In Guillermo A. Pérez and Jean-François Raskin, editors, *34th International Conference on Concurrency Theory, CONCUR 2023, Antwerp, Belgium, September 18-23, 2023*, volume 279 of *LIPIcs*, pages 39:1–39:18. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2023. doi:10.4230/LIPIcs.CONCUR.2023.39.
- 8 Mohamed Faouzi Atig, Ahmed Bouajjani, and Shaz Qadeer. Context-bounded analysis for concurrent programs with dynamic creation of threads. *Log. Methods Comput. Sci.*, 7(4), 2011. doi:10.2168/LMCS-7(4:4)2011.
- 9 Mohamed Faouzi Atig, Ahmed Bouajjani, and Tayssir Touili. On the reachability analysis of acyclic networks of pushdown systems. In Franck van Breugel and Marsha Chechik, editors, *CONCUR 2008 - Concurrency Theory, 19th International Conference, CONCUR 2008, Toronto, Canada, August 19-22, 2008. Proceedings*, volume 5201 of *Lecture Notes in Computer Science*, pages 356–371. Springer, 2008. doi:10.1007/978-3-540-85361-9\_29.
- 10 Mohamed Faouzi Atig, Dmitry Chistikov, Piotr Hofman, K. Narayan Kumar, Prakash Saivasan, and Georg Zetsche. The complexity of regular abstractions of one-counter languages. In Martin Grohe, Eric Koskinen, and Natarajan Shankar, editors, *Proceedings of the 31st Annual ACM/IEEE Symposium on Logic in Computer Science, LICS '16, New York, NY, USA, July 5-8, 2016*, pages 207–216. ACM, 2016. doi:10.1145/2933575.2934561.
- 11 Mohamed Faouzi Atig, Roland Meyer, Sebastian Muskalla, and Prakash Saivasan. On the upward/downward closures of Petri nets. In Kim G. Larsen, Hans L. Bodlaender, and Jean-François Raskin, editors, *42nd International Symposium on Mathematical Foundations of Computer Science, MFCS 2017, Aalborg, Denmark, August 21-25, 2017*, volume 83 of *LIPIcs*, pages 49:1–49:14. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2017. doi:10.4230/LIPIcs.MFCS.2017.49.
- 12 Georg Bachmeier, Michael Luttenberger, and Maximilian Schlund. Finite automata for the sub- and superword closure of CFLs: Descriptive and computational complexity. In Adrian-Horia Dediu, Enrico Formenti, Carlos Martín-Vide, and Bianca Truthe, editors, *Language and Automata Theory and Applications - 9th International Conference, LATA 2015, Nice, France, March 2-6, 2015, Proceedings*, volume 8977 of *Lecture Notes in Computer Science*, pages 473–485. Springer, 2015. doi:10.1007/978-3-319-15579-1\_37.
- 13 David Barozzini, Lorenzo Clemente, Thomas Colcombet, and Pawel Parys. Cost automata, safe schemes, and downward closures. *Fundam. Informaticae*, 188(3):127–178, 2022. doi:10.3233/FI-222145.
- 14 David Barozzini, Pawel Parys, and Jan Wroblewski. Unboundedness for recursion schemes: A simpler type system. In Mikolaj Bojanczyk, Emanuela Merelli, and David P. Woodruff, editors, *49th International Colloquium on Automata, Languages, and Programming, ICALP 2022, Paris, France, July 4-8, 2022*, volume 229 of *LIPIcs*, pages 112:1–112:19. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2022. doi:10.4230/LIPIcs.ICALP.2022.112.
- 15 Pascal Baumann, Moses Ganardi, Rupak Majumdar, Ramanathan S. Thinniyam, and Georg Zetsche. Context-bounded analysis of concurrent programs (invited talk). In Kousha Etesami, Uriel Feige, and Gabriele Puppis, editors, *50th International Colloquium on Automata, Languages, and Programming, ICALP 2023, Paderborn, Germany, July 10-14, 2023*, volume 261 of *LIPIcs*, pages 3:1–3:16. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2023. doi:10.4230/LIPIcs.ICALP.2023.3.
- 16 Pascal Baumann, Moses Ganardi, Rupak Majumdar, Ramanathan S. Thinniyam, and Georg Zetsche. Checking refinement of asynchronous programs against context-free specifications. In Kousha Etesami, Uriel Feige, and Gabriele Puppis, editors, *50th International Colloquium on Automata, Languages, and Programming, ICALP 2023, Paderborn, Germany, July 10-14, 2023*, volume 261 of *LIPIcs*, pages 110:1–110:20. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2023. doi:10.4230/LIPIcs.ICALP.2023.110.

- 17 Pascal Baumann, Rupak Majumdar, Ramanathan S. Thinniyam, and Georg Zetsche. Context-bounded verification of liveness properties for multithreaded shared-memory programs. *Proc. ACM Program. Lang.*, 5(POPL):1–31, 2021. doi:10.1145/3434325.
- 18 Pascal Baumann, Rupak Majumdar, Ramanathan S. Thinniyam, and Georg Zetsche. Context-bounded verification of thread pools. *Proc. ACM Program. Lang.*, 6(POPL):1–28, 2022. doi:10.1145/3498678.
- 19 Noam Chomsky. On certain formal properties of grammars. *Inf. Control.*, 2(2):137–167, 1959. doi:10.1016/S0019-9958(59)90362-6.
- 20 Lorenzo Clemente, Pawel Parys, Sylvain Salvati, and Igor Walukiewicz. The diagonal problem for higher-order recursion schemes is decidable. In Martin Grohe, Eric Koskinen, and Natarajan Shankar, editors, *Proceedings of the 31st Annual ACM/IEEE Symposium on Logic in Computer Science, LICS '16, New York, NY, USA, July 5-8, 2016*, pages 96–105. ACM, 2016. doi:10.1145/2933575.2934527.
- 21 Thomas Colcombet. A combinatorial theorem for trees. In Lars Arge, Christian Cachin, Tomasz Jurdziński, and Andrzej Tarlecki, editors, *Automata, Languages and Programming*, pages 901–912. Springer, 2007. doi:10.1007/978-3-540-73420-8\_77.
- 22 Bruno Courcelle. On constructing obstruction sets of words. *Bull. EATCS*, 44:178–186, 1991.
- 23 Bruno Courcelle and Géraud Sénizergues. The obstructions of a minor-closed set of graphs defined by hyperedge replacement can be constructed. In Janice E. Cuny, Hartmut Ehrig, Gregor Engels, and Grzegorz Rozenberg, editors, *Graph Grammars and Their Application to Computer Science, 5th International Workshop, Williamsburg, VA, USA, November 13-18, 1994, Selected Papers*, volume 1073 of *Lecture Notes in Computer Science*, pages 351–367. Springer, 1994. doi:10.1007/3-540-61228-9\_98.
- 24 Joost Engelfriet. Iterated stack automata and complexity classes. *Inf. Comput.*, 95(1):21–75, 1991. doi:10.1016/0890-5401(91)90015-T.
- 25 Szilárd Zsolt Fazekas, Béla Klein, Tore Koß, Florin Manea, Robert Mercas, and Timo Specht. Subsequence matching and analysis problems for automata with translucent letters. In Giuseppa Castiglione and Sabrina Mantaci, editors, *Implementation and Application of Automata - 29th International Conference, CIAA 2025, Palermo, Italy, September 22-25, 2025, Proceedings*, Lecture Notes in Computer Science, pages 148–164. Springer, 2025. doi:10.1007/978-3-032-02602-6\_11.
- 26 Szilárd Zsolt Fazekas, Tore Koß, Florin Manea, Robert Mercas, and Timo Specht. Subsequence Matching and Analysis Problems for Formal Languages. In Julián Mestre and Anthony Wirth, editors, *35th International Symposium on Algorithms and Computation (ISAAC 2024)*, volume 322 of *Leibniz International Proceedings in Informatics (LIPIcs)*, pages 28:1–28:23, Dagstuhl, Germany, 2024. Schloss Dagstuhl – Leibniz-Zentrum für Informatik. doi:10.4230/LIPIcs.ISAAC.2024.28.
- 27 Moses Ganardi, Irmak Sağlam, and Georg Zetsche. Directed regular and context-free languages. In Olaf Beyersdorff, Mamadou Moustapha Kanté, Orna Kupferman, and Daniel Lokshtanov, editors, *41st International Symposium on Theoretical Aspects of Computer Science, STACS 2024, Clermont-Ferrand, France, March 12-14, 2024*, volume 289 of *LIPIcs*, pages 36:1–36:20. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2024. doi:10.4230/LIPIcs.STACS.2024.36.
- 28 Hugo Gimbert, Corto Mascle, and Patrick Totzke. Optimal sequential flows. *arXiv pre-print arXiv:2511.13806*, 2025. To appear in Automata, Languages and Programming: 53rd International Colloquium, ICALP 2026. doi:10.48550/arXiv.2511.13806.
- 29 Hermann Gruber and Markus Holzer. On pumping constants and smallest grammars for context-free languages. In María Dolores Jiménez-López and György Vaszil, editors, *Languages of Cooperation and Communication - Essays Dedicated to Erzsébet Csuhaj-Varjú to Celebrate Her Scientific Career*, Lecture Notes in Computer Science, pages 54–68. Springer, 2025. doi:10.1007/978-3-031-97274-4\_4.

- 30 Hermann Gruber, Markus Holzer, and Martin Kutrib. More on the size of Higman-Haines sets: Effective constructions. *Fundam. Informaticae*, 91(1):105–121, 2009. doi:10.3233/FI-2009-0035.
- 31 Peter Habermehl, Roland Meyer, and Harro Wimmel. The downward-closure of Petri net languages. In Samson Abramsky, Cyril Gavoille, Claude Kirchner, Friedhelm Meyer auf der Heide, and Paul G. Spirakis, editors, *Automata, Languages and Programming, 37th International Colloquium, ICALP 2010, Bordeaux, France, July 6-10, 2010, Proceedings, Part II*, volume 6199 of *Lecture Notes in Computer Science*, pages 466–477. Springer, 2010. doi:10.1007/978-3-642-14162-1\_39.
- 32 Matthew Hague, Jonathan Kochems, and C.-H. Luke Ong. Unboundedness and downward closures of higher-order pushdown automata. In Rastislav Bodík and Rupak Majumdar, editors, *Proceedings of the 43rd Annual ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages, POPL 2016, St. Petersburg, FL, USA, January 20 - 22, 2016*, pages 151–163. ACM, 2016. doi:10.1145/2837614.2837627.
- 33 Leonard H. Haines. On free monoids partially ordered by embedding. *Journal of Combinatorial Theory*, 6(1):94–98, 1969. doi:10.1016/S0021-9800(69)80111-0.
- 34 Takeshi Hayashi. On derivation trees of indexed grammars—an extension of the uvwxy-theorem—. *Publications of the Research Institute for Mathematical Sciences*, 9(1):61–92, 1973. doi:10.2977/PRIMS/1195192738.
- 35 Graham Higman. Ordering by divisibility in abstract algebras. *Proceedings of the London Mathematical Society*, s3-2(1):326–336, 1952. doi:10.1112/plms/s3-2.1.326.
- 36 Markus Holzer and Christian Rauch. On minimal pumping constants for regular languages. In Zolt Gazdag, Szabolcs Iván, and Gergely Kovásznai, editors, *Proceedings of the 16th International Conference on Automata and Formal Languages, AFL 2023, Eger, Hungary, September 5-7, 2023*, EPTCS, pages 127–141, 2023. doi:10.4204/EPTCS.386.11.
- 37 Ismaël Jecker. A Ramsey Theorem for Finite Monoids. In Markus Bläser and Benjamin Monmege, editors, *38th International Symposium on Theoretical Aspects of Computer Science, STACS 2021, March 16-19, 2021, Saarbrücken, Germany (Virtual Conference)*, volume 187 of *LIPIcs*, pages 44:1–44:13. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2021. doi:10.4230/LIPIcs.STACS.2021.44.
- 38 Ismaël Jecker. A Ramsey theorem for finite monoids. *CoRR*, abs/2101.05895, 2021. doi:10.48550/arXiv.2101.05895.
- 39 Prateek Karandikar, Matthias Niewerth, and Philippe Schnoebelen. On the state complexity of closures and interiors of regular languages with subwords and superwords. *Theor. Comput. Sci.*, 610:91–107, 2016. doi:10.1016/J.TCS.2015.09.028.
- 40 Alexander Kartzow. A pumping lemma for collapsible pushdown graphs of level 2. In Marc Bezem, editor, *Computer Science Logic - 25th International Workshop / 20th Annual Conference of the EACSL, CSL 2011, Bergen, Norway, September 12-15, 2011, Proceedings*, volume 12 of *LIPIcs*, pages 322–336. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2011. doi:10.4230/LIPIcs.CSL.2011.322.
- 41 Teodor Knapik, Damian Niwinski, and Pawel Urzyczyn. Higher-order pushdown trees are easy. In Mogens Nielsen and Uffe Engberg, editors, *Foundations of Software Science and Computation Structures, 5th International Conference, FOSSACS 2002. Held as Part of the Joint European Conferences on Theory and Practice of Software, ETAPS 2002 Grenoble, France, April 8-12, 2002, Proceedings*, volume 2303 of *Lecture Notes in Computer Science*, pages 205–222. Springer, 2002. doi:10.1007/3-540-45931-6\_15.
- 42 Naoki Kobayashi. Types and higher-order recursion schemes for verification of higher-order programs. In Zhong Shao and Benjamin C. Pierce, editors, *Proceedings of the 36th ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages, POPL 2009, Savannah, GA, USA, January 21-23, 2009*, pages 416–428. ACM, 2009. doi:10.1145/1480881.1480933.
- 43 Naoki Kobayashi. Model checking higher-order programs. *J. ACM*, 60(3):20:1–20:62, 2013. doi:10.1145/2487241.2487246.



- 44 Salvatore La Torre, Anca Muscholl, and Igor Walukiewicz. Safety of parametrized asynchronous shared-memory systems is almost always decidable. In Luca Aceto and David de Frutos-Escrig, editors, *26th International Conference on Concurrency Theory, CONCUR 2015, Madrid, Spain, September 1-4, 2015*, volume 42 of *LIPIcs*, pages 72–84. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2015. doi:10.4230/LIPIcs.CONCUR.2015.72.
- 45 Théodore Lopez, Benjamin Monmege, and Jean-Marc Talbot. Regular d-length: A tool for improved prefix-stable forward Ramsey factorisations. *Inf. Process. Lett.*, 187:106497, 2025. doi:10.1016/J.IPL.2024.106497.
- 46 Rupak Majumdar, Ramanathan S. Thinniyam, and Georg Zetsche. General decidability results for asynchronous shared-memory programs: Higher-order and beyond. *Log. Methods Comput. Sci.*, 18(4), 2022. doi:10.46298/LMCS-18(4:2)2022.
- 47 Richard Mandel, Corto Mascle, and Georg Zetsche. The complexity of downward closures of indexed languages. *CoRR*, abs/2601.19466, 2026. doi:10.48550/arXiv.2601.19466.
- 48 Richard Mayr. Undecidable problems in unreliable computations. *Theor. Comput. Sci.*, 297(1-3):337–354, 2003. doi:10.1016/S0304-3975(02)00646-1.
- 49 Luke Ong. Higher-order model checking: An overview. In *30th Annual ACM/IEEE Symposium on Logic in Computer Science, LICS 2015, Kyoto, Japan, July 6-10, 2015*, pages 1–15. IEEE Computer Society, 2015. doi:10.1109/LICS.2015.9.
- 50 Pawel Parys. The complexity of the diagonal problem for recursion schemes. In Satya V. Lokam and R. Ramanujam, editors, *37th IARCS Annual Conference on Foundations of Software Technology and Theoretical Computer Science, FSTTCS 2017, Kanpur, India, December 11-15, 2017*, volume 93 of *LIPIcs*, pages 45:1–45:14. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2017. doi:10.4230/LIPIcs.FSTTCS.2017.45.
- 51 Pawel Parys. Intersection types for unboundedness problems. In Michele Pagani and Sandra Alves, editors, *Proceedings Twelfth Workshop on Developments in Computational Models and Ninth Workshop on Intersection Types and Related Systems, DCM/ITRS 2018, Oxford, UK, 8th July 2018*, volume 293 of *EPTCS*, pages 7–27, 2018. doi:10.4204/EPTCS.293.2.
- 52 Jacques Sakarovitch. *Elements of Automata Theory*. Cambridge University Press, 2009. URL: <http://www.cambridge.org/uk/catalogue/catalogue.asp?isbn=9780521844253>.
- 53 Sylvain Schmitz. *Algorithmic Complexity of Well-Quasi-Orders*. Accreditation to supervise research, École normale supérieure Paris-Saclay, November 2017. URL: <https://theses.hal.science/tel-01663266>.
- 54 Sylvain Schmitz and Philippe Schnoebelen. Multiply-recursive upper bounds with higman’s lemma. In Luca Aceto, Monika Henzinger, and Jiri Sgall, editors, *Automata, Languages and Programming - 38th International Colloquium, ICALP 2011, Zurich, Switzerland, July 4-8, 2011, Proceedings, Part II*, volume 6756 of *Lecture Notes in Computer Science*, pages 441–452. Springer, 2011. doi:10.1007/978-3-642-22012-8\_35.
- 55 Imre Simon. Factorization forests of finite height. *Theor. Comput. Sci.*, 72(1):65–94, 1990. doi:10.1016/0304-3975(90)90047-L.
- 56 Tim Smith. A new pumping lemma for indexed languages, with an application to infinite words. *Inf. Comput.*, 252:176–186, 2017. doi:10.1016/J.IC.2016.11.002.
- 57 Jan van Leeuwen. Effective constructions in well-partially-ordered free monoids. *Discrete Mathematics*, 21(3):237–252, 1978. doi:10.1016/0012-365X(78)90156-5.
- 58 Georg Zetsche. Computing downward closures for stacked counter automata. In Ernst W. Mayr and Nicolas Ollinger, editors, *32nd International Symposium on Theoretical Aspects of Computer Science, STACS 2015, Garching, Germany, March 4-7, 2015*, volume 30 of *LIPIcs*, pages 743–756. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2015. doi:10.4230/LIPIcs.STACS.2015.743.
- 59 Georg Zetsche. An approach to computing downward closures. In Magnús M. Halldórsson, Kazuo Iwama, Naoki Kobayashi, and Bettina Speckmann, editors, *Automata, Languages, and Programming - 42nd International Colloquium, ICALP 2015, Kyoto, Japan, July 6-10, 2015, Proceedings, Part II*, volume 9135 of *Lecture Notes in Computer Science*, pages 440–451. Springer, 2015. doi:10.1007/978-3-662-47666-6\_35.

- 60   Georg Zetsche. The complexity of downward closure comparisons. In Ioannis Chatzigiannakis, Michael Mitzenmacher, Yuval Rabani, and Davide Sangiorgi, editors, *43rd International Colloquium on Automata, Languages, and Programming, ICALP 2016, Rome, Italy, July 11-15, 2016*, volume 55 of *LIPIcs*, pages 123:1–123:14. Schloss Dagstuhl – Leibniz-Zentrum für Informatik, 2016. doi:10.4230/LIPIcs.ICALP.2016.123.
- 61   Georg Zetsche. Separability by piecewise testable languages and downward closures beyond subwords. In Anuj Dawar and Erich Grädel, editors, *Proceedings of the 33rd Annual ACM/IEEE Symposium on Logic in Computer Science, LICS 2018, Oxford, UK, July 09-12, 2018*, pages 929–938. ACM, 2018. doi:10.1145/3209108.3209201.