



PONTIFICIA UNIVERSIDAD CATÓLICA DE CHILE
ESCUELA DE INGENIERÍA

EXPRESSIVE POWER OF LINEAR ALGEBRA QUERY LANGUAGES

THOMAS MUÑOZ SERRANO

Thesis submitted to the Office of Research and Graduate Studies
in partial fulfillment of the requirements for the degree of
Master of Science in Engineering

Advisors:

CRISTIAN RIVEROS, DOMAGOJ VRGOČ

Santiago de Chile, October 2020

© 2020, THOMAS MUÑOZ SERRANO



PONTIFICIA UNIVERSIDAD CATÓLICA DE CHILE
ESCUELA DE INGENIERÍA

EXPRESSIVE POWER OF LINEAR ALGEBRA QUERY LANGUAGES

THOMAS MUÑOZ SERRANO

Members of the Committee:

CRISTIAN RIVEROS

DOMAGOJ VRGOČ

PABLO BARCELÓ

JORGE PÉREZ

LUIS RIZZI

Domagoj Vrgoč

Thesis submitted to the Office of Research and Graduate Studies
in partial fulfillment of the requirements for the degree of
Master of Science in Engineering

Santiago de Chile, October 2020

© 2020, THOMAS MUÑOZ SERRANO

ACKNOWLEDGEMENTS

My loving parents, Luis and María Isabel, thank you for raising me in so many ways, for every teaching they gave me, for every advice, distraction, comfort, love and support. I would not be the person I am today without you, thank you so much.

My brothers, Lukas and David, thank you for being there always, despite you almost never mess with my academic stuff, we joke around and express our love and support in multiple ways. Thanks you for all the games we played, the matches we had and the constant kidding around.

My girlfriend, Evelyn, I met you shortly before I started my masters and you are the most amazing person, your love and support meant a lot during this time. Thank you for all the advice, love and laughs.

My grandparents, uncles, aunts and cousins, for being there all my life and show nothing but love and support all the time.

My friends, all of you: the ones from childhood, the ones from school, the ones from college, the ones from sports; the ones from life. We haven't had the chance to see each other as much as I would like, but keep in mind that every single one of you are always in my thoughts, each one of you has made a huge difference in me and I have learned something from everyone. Thank you for all the jokes, the conversations, the discussions, the love, the hate, the good and the bad moments, everything.

My teachers from school, thank you all. You taught so many human values and how they relate to academic work. I will always be grateful for that.

Cristian Riveros, one of my teachers one of my supervisors, and my friend. Thank you for everything that you taught me, for the potential you saw in me and for giving me the opportunity to do academic work and be part of your teaching team. It has been an

amazing experience and I cannot thank you enough. It was an honor to work with you and I look forward to do more.

Domagoj Vrgoč, also one of my teachers and one of my supervisors, and my friend. Thank you for always treating me as another academic and for always being available when I needed to discuss ideas. Thank you for everything you taught me, I learned so much from you. It has been a pleasure to work with you and look forward to do more.

To all the people that make modern gaming possible. Thank you for all the enjoyment and distractions, crucial for my mental health.

Finally, I would like to thank myself very much, for all the efforts I made, the problems I overcame, the help I asked, the moments I enjoyed, the love I gave and the work I made. I accomplished what I wanted for many years now: to contribute science to the world. I made it, I'm a scientist.

TABLE OF CONTENTS

ACKNOWLEDGEMENTS	iii
LIST OF FIGURES	vii
ABSTRACT	viii
RESUMEN	ix
1. INTRODUCTION	1
1.1. Contribution and outline	4
1.2. Related work	5
2. MATLANG	6
2.1. Syntax	6
2.2. Schemas and typing	7
2.3. Semantics	7
3. Extending MATLANG with for loops	10
3.1. Syntax and semantics of <code>for</code> -MATLANG	10
3.2. Design decisions behind <code>for</code> -MATLANG	13
3.2.1. Loop Initialization	13
3.2.2. Order	15
3.3. Syntactic definitions and simplifications	19
3.3.1. Operators	19
3.3.2. Function application simplifications	21
4. Algorithms in Linear Algebra	22
4.1. LU decomposition	22
4.1.1. LU decomposition by Gaussian elimination	22
4.1.2. Implementing LU decomposition in <code>for</code> -MATLANG	23
4.2. Determinant and inverse	28

5. Expressiveness of for loops	34
5.1. Linear space functions	34
5.2. From arithmetic circuits to <code>for-MATLANG</code>	47
5.3. From <code>for-MATLANG</code> to circuits	62
5.4. Supporting additional operators	70
6. Restricting the power of for loops	72
6.1. Summation matlang and relational algebra	72
6.2. Hadamard product and weighted logics	84
6.3. Matrix multiplication as a quantifier	91
7. CONCLUSIONS	94
REFERENCES	96

LIST OF FIGURES

6.1	Fragments of <code>for-MATLANG</code> and their equivalences. The functions <code>4CLIQUE</code> , <code>DP</code> (diagonal product), <code>INV</code> , <code>DET</code> , and <code>PLU</code> decomposition are placed in their fragments.	93
-----	--	----

ABSTRACT

Linear algebra algorithms often require some sort of iteration or recursion as is illustrated by standard algorithms for Gaussian elimination, matrix inversion, and transitive closure. A key characteristic shared by these algorithms is that they allow looping for a number of steps that is bounded by the matrix dimension. In this thesis we extend the matrix query language `MATLANG` with this type of recursion, and show that this suffices to express classical linear algebra algorithms. We study the expressive power of this language and show that it naturally corresponds to arithmetic circuit families, which are often said to capture linear algebra. Furthermore, we analyze several sub-fragments of our language, and show that their expressive power is closely tied to logical formalisms on semiring-annotated relations.

Keywords: logic, matlang, linear algebra, arithmetic circuits, relational algebra.

RESUMEN

Los algoritmos del álgebra lineal a menudo requieren algún tipo de iteración o recursión, como lo ilustran los algoritmos estándar para la eliminación gaussiana, la inversión de matrices y la clausura transitiva. Una característica clave compartida por estos algoritmos es que permiten que se repitan varios pasos, pero limitados por la dimensión de la matriz. En esta tesis, ampliamos el lenguaje de consulta para matrices `MATLANG` con este tipo de recursión, y evidenciamos que esto es suficiente para expresar algoritmos clásicos del álgebra lineal. Estudiamos el poder expresivo de este lenguaje y demostramos que corresponde naturalmente a las familias de circuitos aritméticos, que a menudo se dice que capturan el álgebra lineal. Además, analizamos varios sub-fragmentos de nuestro lenguaje, y se demuestra que su poder expresivo está estrechamente relacionado con formalismos lógicos en las relaciones anotadas en semi-anillos.

Palabras Claves: lógica, matlang, álgebra lineal, circuitos aritméticos, álgebra relacional.

1. INTRODUCTION

Linear algebra-based algorithms have become a key component in data analytic workflows. As such, there is a growing interest in the database community to integrate linear algebra functionalities into relational database management systems (Luo, Gao, Gubanov, Perez, & Jermaine, 2018; Boehm, Kumar, & Yang, 2019; Kunft, Alexandrov, Katsifodimos, & Markl, 2016; Jankov et al., 2019; Khamis, Ngo, Nguyen, Olteanu, & Schleich, 2018). In particular, from a query language perspective, several proposals have recently been put forward to unify relational algebra and linear algebra. Two notable examples of this are `LARA` (Hutchison, Howe, & Suciu, 2017), a minimalistic language in which a number of atomic operations on associative tables are proposed, and `MATLANG`, a query language for matrices (Brijder, Geerts, den Bussche, & Weerwag, 2019).

Both `LARA` and `MATLANG` have been studied by the database theory community, showing interesting connections to relational algebra and logic. For example, fragments of `LARA` are known to capture first-order logic with aggregation (Barceló, Higuera, Pérez, & Subercaseaux, 2020), and `MATLANG` has been recently shown to be equivalent to a restricted version of the (positive) relational algebra on K -relations, RA_K^+ (Brijder, Gyssens, & den Bussche, 2020), where K denotes a semiring. On the other hand, some standard constructions in linear algebra are out of reach for these languages. For instance, it was shown that under standard complexity-theoretic assumptions, `LARA` can not compute the inverse of a matrix or its determinant (Barceló et al., 2020), and operations such as the transitive closure of a matrix are known to be inexpressible in `MATLANG` (Brijder et al., 2019). Given that these are fundamental constructs in linear algebra, one might wonder how to extend `LARA` or `MATLANG` in order to allow expressing such properties.

One approach would be to add these constructions explicitly to the language. Indeed, this was done for `MATLANG` by Brijder et al., 2019. In this work, the authors have extended the core language with the trace, the inverse, the determinant, or the eigenvectors operators and study the expressive power of the result. However, one can argue that there is nothing

special in these operators, apart that they have been used historically in linear algebra textbooks and they extend the expressibility of the core language. The question here is whether these new operators form a sound and natural choice to extend the core language, or whether they are just some particular queries that we would like to support. A similar approach is done with LARA by Barceló et al., 2020, where instead of explicitly adding the convolution operator, they extend the language with a lineal order over the domain, and prove that, as a consequence, convolution can be computed.

In this thesis we take the more principled approach by studying what are the atomic operations needed to define standard linear algebra algorithms. Inspecting any linear algebra textbook, one sees that most linear algebra procedures heavily rely on the use of for-loops in which iterations happen over the dimensions of the matrices involved. To illustrate this, let us consider the example of computing the transitive closure of a graph. This can be done using a modification of the Floyd-Warshall algorithm (Cormen, Leiserson, Rivest, & Stein, 2001), which takes as its input an $n \times n$ adjacency matrix A representing our graph, and operates according to the following pseudo-code:

```

for  $k = 1..n$  do
  for  $i = 1..n$  do
    for  $j = 1..n$  do
       $A[i, j] := A[i, j] + A[i, k] \cdot A[k, j]$ 

```

After executing the algorithm, all of the non zero entries represent the existence of an edge in the (irreflexive) transitive closure graph.

By examining standard linear algebra algorithms such as Gaussian elimination, LU -decomposition, computing the inverse of a matrix, or its determinant, we can readily see that this pattern continues. Namely, we observe that there are two main components to such algorithms: (i) the ability to iterate up to the matrix dimension; and (ii) the ability to access a particular position in our matrix. In order to allow this behavior in a query language, we propose to extend MATLANG with limited recursion in the form of for-loops,

resulting in the language `for-MATLANG`. To simulate the two components of standard linear algebra algorithms in a natural way, we simulate a loop of the form `for` $i = 1..n$ `do` by leveraging canonical vectors. In other words, we use the canonical vectors $b_1 = (1, 0, \dots)$, $b_2 = (0, 1, \dots)$, $\dots$, to access specific rows and columns, and iterate over these vectors. In this way, we obtain a language able to compute important linear algebra operators such as LU -decomposition, determinant, matrix inverse, among other things.

Of course, a natural question to ask now is whether this really results in a language suitable for linear algebra? We argue that the correct way to approach this question is to compare our language to arithmetic circuits, which have been shown to capture the vast majority of existing matrix algorithms, from basic ones such as computing the determinant and the inverse, to complex procedures such as discrete Fourier transformation, and Strassen’s algorithm (see Shpilka & Yehudayoff, 2010; Allender, 2004 for an overview of the area), and can therefore be considered to effectively capture linear algebra. In the main technical result of this thesis, we show that `for-MATLANG` indeed computes the same class of functions over matrices as the ones computed by arithmetic circuit families of bounded degree. As a consequence, `for-MATLANG` inherits all expressiveness properties of circuits, and thus can simulate any linear algebra algorithm definable by circuits.

Having established that `for-MATLANG` indeed provides a good basis for a linear algebra language, we move to a more fine-grained analysis of the expressiveness of its different fragments. For this, we aim to provide a connection with logical formalisms, similarly as was done by linking `LARA` and `MATLANG` to first-order logic with aggregates (Barceló et al., 2020; Brijder et al., 2019). As we show, capturing different logics correspond to restricting how matrix variables are updated in each iteration of the `for`-loops allowed in `for-MATLANG`. For instance, if we only allow to add some temporary result to a variable in each iteration (instead of rewriting it completely like in any programming language), we obtain a language, called `sum-MATLANG`, which is equivalent to RA_K^+ , directly extending an analogous result shown for `MATLANG`, mentioned earlier (Brijder et al., 2020). We then study updating matrix variables based on another standard linear algebra operator,

the Hadamard product, resulting in a fragment called `FO-MATLANG`, which we show to be equivalent to weighted logics (Droste & Gastin, 2005). Finally, in `prod-MATLANG` we consider updating the variables based on the standard matrix product, and link this fragment of our language to the ones discussed previously.

1.1. Contribution and outline

After we recall `MATLANG` in Chapter 2, we show in Chapter 3 how for-loops can be added to `MATLANG` in a natural way. We also observe that `for-MATLANG` strictly extends `MATLANG`. In addition, we discuss some design decisions behind the definition of `for-MATLANG`, noting that our use of canonical vectors results in the availability of an order relation.

In Chapter 4 we show that `for-MATLANG` can compute important linear algebra algorithms in a natural way. We provide expressions in `for-MATLANG` for LU decomposition (used to solve linear systems of equations), the determinant and matrix inversion.

More generally, in Chapter 5 we report our main technical contribution. We show that every uniform arithmetic circuits of polynomial degree correspond to a `for-MATLANG` expression, and vice versa, when a `for-MATLANG` expression has polynomial degree, then there is an equivalent uniform family of arithmetic circuits. As a consequence, `for-MATLANG` inherits all expressiveness properties of such circuits.

Finally, in Chapter 6 we generalize the semantics of `for-MATLANG` to matrices with values in a semiring K , and show that two natural fragment of `for-MATLANG`, `sum-MATLANG`, and `FO-MATLANG`, are equivalent to the (positive) relational algebra and weighted logics on binary K -relations, respectively. We also briefly comment on a minimal fragment of `for-MATLANG`, based on `prod-MATLANG`, that is able to compute matrix inversion.

1.2. Related work

We already mentioned `LARA` (Hutchison et al., 2017) and `MATLANG` (Brijder et al., 2019) whose expressive power was further analyzed by Barceló et al., 2020; Brijder et al., 2020; Geerts, 2019, 2020. Extensions of `SQL` for matrix manipulations are reported by Luo et al., 2018. Most relevant is the work of Jankov et al., 2019 in which a recursion mechanism is added to `SQL` which resembles for-loops. The expressive power of this extension is unknown, however. Classical logics with aggregation (Hella, Libkin, Nurmonen, & Wong, 2001) and fixed-point logics with counting (Grohe & Pakusa, 2017) can also be used for linear algebra. More generally, for the descriptive complexity of linear algebra we refer to the work by Dawar, Grohe, Holm, & Laubner, 2009; Holm, 2010. Most of these works require to encode real numbers inside relations, whereas we treat real numbers as atomic values. We refer to relevant papers related to arithmetic circuits and logical formalisms on semiring-annotated relations in the corresponding chapters later in the thesis.

2. MATLANG

We start by recalling the matrix query language MATLANG, introduced by Brijder et al., 2019, which serves as our starting point.

2.1. Syntax

Let $\mathcal{V} = \{V_1, V_2, \dots\}$ be a countably infinite set of *matrix variables* and $\mathcal{F} = \bigcup_{k \geq 1} \mathcal{F}_k$ with $\mathcal{F}_k$ a set of *functions* of the form $f : \mathbb{R}^k \rightarrow \mathbb{R}$, where $\mathbb{R}$ denotes the set of real numbers. The syntax of MATLANG expressions is defined by the following grammar¹:

$e ::=$	$V \in \mathcal{V}$	(matrix variable)
	$ \quad e^T$	(transpose)
	$ \quad \mathbf{1}(e)$	(one-vector)
	$ \quad \text{diag}(e)$	(diagonalization of a vector)
	$ \quad e_1 \cdot e_2$	(matrix multiplication)
	$ \quad e_1 + e_2$	(matrix addition)
	$ \quad e_1 \times e_2$	(scalar multiplication)
	$ \quad f(e_1, \dots, e_k)$	(pointwise application of $f \in \mathcal{F}_k$).

MATLANG is parametrized by a collection of functions $\mathcal{F}$ but in the remainder of the thesis we only make this dependence explicit, and write $\text{MATLANG}[\mathcal{F}]$, for some set $\mathcal{F}$ of functions, when these functions are crucial for some results to hold. When we simply write MATLANG, we mean that any function can be used (including not using any function at all).

¹The original syntax also permits the operator $\text{let } V = e_1 \text{ in } e_2$, which replaces every occurrence of V in e_2 with the value of e_1 . Since this is just syntactic sugar, we omit this operator. We also explicitly include matrix addition and scalar multiplication, although these can be simulated by pointwise function applications. Finally, we use transposition instead of conjugate transposition since we work with matrices over $\mathbb{R}$.

2.2. Schemas and typing

To define the semantics of MATLANG expressions we need a notion of schema and well-typedness of expressions. A MATLANG *schema* $\mathcal{S}$ is a pair $\mathcal{S} = (\mathcal{M}, \text{size})$, where $\mathcal{M} \subset \mathcal{V}$ is a finite set of matrix variables, and $\text{size} : \mathcal{M} \mapsto \text{Symb} \times \text{Symb}$ is a function that maps each matrix variable in $\mathcal{M}$ to a pair of *size symbols*. The size function helps us determine whether certain matrix operations, such as matrix multiplication, can be performed for matrices adhering to a schema. We denote size symbols by Greek letters α, β, γ . We also assume that $1 \in \text{Symb}$. To help us determine whether a MATLANG expression can always be evaluated, we define the *type* of an expression e , with respect to a schema $\mathcal{S}$, denoted by $\text{type}_{\mathcal{S}}(e)$, inductively as follows:

- $\text{type}_{\mathcal{S}}(V) := \text{size}(V)$, for a matrix variable $V \in \mathcal{M}$;
- $\text{type}_{\mathcal{S}}(e^T) := (\beta, \alpha)$ if $\text{type}_{\mathcal{S}}(e) = (\alpha, \beta)$;
- $\text{type}_{\mathcal{S}}(\mathbf{1}(e)) := (\alpha, 1)$ if $\text{type}_{\mathcal{S}}(e) = (\alpha, \beta)$;
- $\text{type}_{\mathcal{S}}(\text{diag}(e)) := (\alpha, \alpha)$, if $\text{type}_{\mathcal{S}}(e) = (\alpha, 1)$;
- $\text{type}_{\mathcal{S}}(e_1 \cdot e_2) := (\alpha, \gamma)$ if $\text{type}_{\mathcal{S}}(e_1) = (\alpha, \beta)$, and $\text{type}_{\mathcal{S}}(e_2) = (\beta, \gamma)$;
- $\text{type}_{\mathcal{S}}(e_1 + e_2) := (\alpha, \beta)$ if $\text{type}_{\mathcal{S}}(e_1) = \text{type}_{\mathcal{S}}(e_2) = (\alpha, \beta)$;
- $\text{type}_{\mathcal{S}}(e_1 \times e_2) := (\alpha, \beta)$ if $\text{type}_{\mathcal{S}}(e_1) = (1, 1)$ and $\text{type}_{\mathcal{S}}(e_2) = (\alpha, \beta)$.
- $\text{type}_{\mathcal{S}}(f(e_1, \dots, e_k)) := (\alpha, \beta)$, whenever $\text{type}_{\mathcal{S}}(e_1) = \dots = \text{type}_{\mathcal{S}}(e_k) := (\alpha, \beta)$ and $f \in \mathcal{F}_k$.

We call an expression *well-typed* according to the schema $\mathcal{S}$, if it has a defined type. A well-typed expression can be evaluated regardless of the actual sizes of the matrices assigned to matrix variables, as we describe next.

2.3. Semantics

We use $\text{Mat}[\mathbb{R}]$ to denote the set of all real matrices and for $A \in \text{Mat}[\mathbb{R}]$, $\dim(A) \in \mathbb{N}^2$ denotes its dimensions. A (MATLANG) *instance* $\mathcal{I}$ over a schema $\mathcal{S}$ is a pair $\mathcal{I} = (\mathcal{D}, \text{mat})$, where $\mathcal{D} : \text{Symb} \mapsto \mathbb{N}$ assigns a value to each size symbol (and thus in turn dimensions to

each matrix variable), and $\text{mat} : \mathcal{M} \mapsto \text{Mat}[\mathbb{R}]$ assigns a concrete matrix to each matrix variable $V \in \mathcal{M}$, such that $\dim(\text{mat}(V)) = \mathcal{D}(\alpha) \times \mathcal{D}(\beta)$ if $\text{size}(V) = (\alpha, \beta)$. That is, an instance tells us the dimensions of each matrix variable, and also the concrete matrices assigned to the variable names in $\mathcal{M}$. We assume that $\mathcal{D}(1) = 1$, for every instance $\mathcal{I}$. If e is a well-typed expression according to $\mathcal{S}$, then we denote by $\llbracket e \rrbracket(\mathcal{I})$ the matrix obtained by evaluating e over $\mathcal{I}$, and define it as follows:

- $\llbracket V \rrbracket(\mathcal{I}) := \text{mat}(V)$, for $V \in \mathcal{M}$;
- $\llbracket e^T \rrbracket(\mathcal{I}) := \llbracket e \rrbracket(\mathcal{I})^T$, where A^T is the transpose of matrix A ;
- $\llbracket \mathbf{1}(e) \rrbracket(\mathcal{I})$ is a $n \times 1$ vector with 1 as all of its entries, where $\dim(\llbracket e \rrbracket(\mathcal{I})) = (n, m)$;
- $\llbracket \text{diag}(e) \rrbracket(\mathcal{I})$ is a diagonal matrix with the vector $\llbracket e \rrbracket(\mathcal{I})$ on its main diagonal, and zero in every other position;
- $\llbracket e_1 \cdot e_2 \rrbracket(\mathcal{I}) := \llbracket e_1 \rrbracket(\mathcal{I}) \cdot \llbracket e_2 \rrbracket(\mathcal{I})$;
- $\llbracket e_1 + e_2 \rrbracket(\mathcal{I}) := \llbracket e_1 \rrbracket(\mathcal{I}) + \llbracket e_2 \rrbracket(\mathcal{I})$;
- $\llbracket e_1 \times e_2 \rrbracket(\mathcal{I}) := a \times \llbracket e_2 \rrbracket(\mathcal{I})$ with $\llbracket e_1 \rrbracket(\mathcal{I}) = [a]$; and
- $\llbracket f(e_1, \dots, e_k) \rrbracket(\mathcal{I})$ is a matrix A of the same size as $\llbracket e_1 \rrbracket(\mathcal{I})$, and where A_{ij} has the value $f(\llbracket e_1 \rrbracket(\mathcal{I})_{ij}, \dots, \llbracket e_k \rrbracket(\mathcal{I})_{ij})$.

This concludes the description and semantics of `MATLANG`. Just as it is commonly done in `MATLANG`, we will sometimes abuse the notation and use 1×1 matrices as scalars and vice versa. For an in-depth discussion on why this is permissible from an expressiveness point of view, see (Brijder et al., 2019).

We next provide some simple example.

Example 2.1. Consider the `MATLANG`[$f_\odot$] expression with $f_\odot : \mathbb{R}^2 \rightarrow \mathbb{R} : (x, y) \mapsto x \cdot y$:

$$\text{cwalk} := (\mathbf{1}(V))^T \cdot f_\odot(V \cdot V, \text{diag}(\mathbf{1}(V)) \cdot \mathbf{1}(V)).$$

Let $\mathcal{S}$ consist of $\mathcal{M} := \{V\}$ and $\text{size}(V) := (\alpha, \alpha)$ such that matrices assigned to V by instances $\mathcal{I}$ over $\mathcal{S}$ are square matrices. It is readily verified that `cwalk` is well-typed and

more specifically, $\text{types}(\text{cwalk}) = (1, 1)$, i.e., it returns an element of $\mathbb{R}$ on any instance $\mathcal{I}$. Let $\mathcal{I}$ be such that $\mathcal{D}(\alpha) = n$ and $\text{mat}(V)$ is an adjacency matrix A of an undirected graph G consisting of n vertices. Then, it is readily verified that $\llbracket \text{cwalk} \rrbracket(\mathcal{I})$ returns the number of paths of length two in G which start in and end at the same vertex. $\square$

Although `MATLANG` forms a solid basis for a matrix query language, it is limited in expressive power. Indeed, `MATLANG` is subsumed by first order logic with aggregates that uses only three variables (Brijder et al., 2019). Hence, no `MATLANG` expression exists that can compute the transitive closure of a graph (represented by its adjacency matrix) or can compute the inverse of a matrix. Rather than extending `MATLANG` with specific linear algebra operators, such as matrix inversion, we next introduce a limited form of recursion in `MATLANG`.

3. EXTENDING MATLANG WITH FOR LOOPS

To extend MATLANG with recursion, we take inspiration from classical linear algebra algorithms, such as those described by Press, Teukolsky, Vetterling, & Flannery, 1992. Many of these algorithms are based on *for-loops* in which the termination condition for each loop is determined by the matrix dimensions. We have seen how the transitive closure of a matrix can be computed using for-loops in the Introduction. Here we add this ability to MATLANG, and show that the resulting language, called `for-MATLANG`, can compute properties outside of the scope of MATLANG. We see more advanced examples, such as Gaussian elimination and matrix inversion, later in the thesis.

3.1. Syntax and semantics of `for-MATLANG`

The syntax of `for-MATLANG` is defined just as for MATLANG but with an extra rule in the grammar:

$$\text{for } v, X . e \quad (\text{canonical for loop, with } v, X \in \mathcal{V}).$$

Intuitively, X is a matrix variable which is iteratively updated according to the expression e . We simulate iterations of the form “`for` $i \in [1..n]$ ” by letting v loop over the *canonical vectors* $b_1^n, \dots, b_n^n$ of dimension n . Here, $b_1^n = [1 \ 0 \ \dots \ 0]^T$, $b_2^n = [0 \ 1 \ 0 \ \dots \ 0]^T$, etc¹. When n is clear from the context we simply write $b_1, b_2, \dots$. In addition, the expression e in the rule above may depend on v .

We next make the semantics precise and start by declaring the type of loop expressions. Given a schema $\mathcal{S}$, the type of a `for-MATLANG` expression e , denoted $\text{type}_{\mathcal{S}}(e)$, is defined inductively as in MATLANG but with following extra rule:

¹It is important that the canonical basis used in the iteration is the standard one (normalized), because we aim to check or get the value of a specific position of a matrix using these vectors. We can always compute these vectors if we have any basis of dimension n .

- $\text{type}_{\mathcal{S}}(\text{for } v, X . e) := (\alpha, \beta)$, if
 $\text{type}_{\mathcal{S}}(e) = \text{type}_{\mathcal{S}}(X) = (\alpha, \beta)$ and $\text{type}_{\mathcal{S}}(v) = (\gamma, 1)$.

We note that $\mathcal{S}$ now necessarily includes v and X as variables and assigns size symbols to them. We also remark that in the definition of the type of $\text{for } v, X . e$, we require that $\text{type}_{\mathcal{S}}(X) = \text{type}_{\mathcal{S}}(e)$ as this expression updates the content of the variable X in each iteration using the result of e . We further restrict the type of v to be a vector, i.e., $\text{type}_{\mathcal{S}}(v) = (\gamma, 1)$, since v will be instantiated with canonical vectors. A `for`-MATLANG expression e is well-typed over a schema $\mathcal{S}$ if its type is defined.

For well-typed expressions we next define their semantics. This is done in an inductive way, just as for `MATLANG`. To define the semantics of $\text{for } v, X . e$ over an instance $\mathcal{I}$, we need the following notation. Let $\mathcal{I}$ be an instance and $V \in \mathcal{M}$. Then $\mathcal{I}[V := A]$ denotes an instance that coincides with $\mathcal{I}$, except that the value of the matrix variable V is given by the matrix A . Assume that $\text{type}_{\mathcal{S}}(v) = (\gamma, 1)$, and $\text{type}_{\mathcal{S}}(e) = (\alpha, \beta)$ and $n := \mathcal{D}(\gamma)$. Then, $\llbracket \text{for } v, X . e \rrbracket(\mathcal{I})$ is defined iteratively, as follows:

- Let $A_0 := \mathbf{0}$ be the zero matrix of size $\mathcal{D}(\alpha) \times \mathcal{D}(\beta)$.
- For $i = 1, \dots, n$, compute $A_i := \llbracket e \rrbracket(\mathcal{I}[v := b_i^n, X := A_{i-1}])$.
- Finally, set $\llbracket \text{for } v, X . e \rrbracket(\mathcal{I}) := A_n$.

For better understanding how `for`-MATLANG works, we next provide some examples. We start by showing that the `one-vector` and `diag` operators are redundant in `for`-MATLANG.

Example 3.1. *We first show how the one-vector operator $\mathbf{1}(e)$ can be expressed using `for` loops. It suffices to consider the expression*

$$e_1 := \text{for } v, X . X + v,$$

with $\text{type}_{\mathcal{S}}(v) = (\alpha, 1) = \text{type}_{\mathcal{S}}(X)$ if $\text{type}_{\mathcal{S}}(e) = (\alpha, \beta)$. This expression is well-typed and is of type $(\alpha, 1)$. When evaluated over some instance $\mathcal{I}$ with $n = \mathcal{D}(\alpha)$, $\llbracket e_1 \rrbracket(\mathcal{I})$ is defined as follows. Initially, $A_0 := \mathbf{0}$. Then $A_i := A_{i-1} + b_i^n$, i.e., the i th canonical

vector is added to A_{i-1} . Finally, $\llbracket e_1 \rrbracket(\mathcal{I}) := A_n$ and this now clearly coincides with $\llbracket 1(e) \rrbracket(\mathcal{I})$. $\square$

Example 3.2. *We next show that the `diag` operator is redundant in `for-MATLANG`. Indeed, it suffices to consider the expression*

$$e_{\text{diag}}(e) := \text{for } v, X. X + (v^T \cdot e) \times v \cdot v^T,$$

where e is a `for-MATLANG` expression of type $(\alpha, 1)$. For this expression to be well-typed, v has to be a vector variable of type $\alpha \times 1$ and X a matrix variable of type (α, α) . Then, $\llbracket e_{\text{diag}} \rrbracket(\mathcal{I})$ is defined as follows. Initially, A_0 is the zero matrix of dimension $n \times n$, where $n = \mathcal{D}(\alpha)$. Then, in each iteration $i \in [1..n]$, $A_i := A_{i-1} + ((b_i^n)^T \cdot \llbracket e \rrbracket(\mathcal{I})) \times (b_i^n \cdot (b_i^n)^T)$. In other words, A_i is obtained by adding the matrix with value $(\llbracket e \rrbracket(\mathcal{I}))_i$ on position (i, i) to A_{i-1} . Hence, $\llbracket e_{\text{diag}} \rrbracket(\mathcal{I}) := A_n = \llbracket \text{diag}(e) \rrbracket(\mathcal{I})$. $\square$

These examples illustrate that we can limit `for-MATLANG` to consist of the following “core” operators: transposition, matrix multiplication and addition, scalar multiplication, pointwise function application, and for-loops. More specific, `for-MATLANG` is defined by the following simplified syntax:

$$e ::= V \mid e^T \mid e_1 \cdot e_2 \mid e_1 + e_2 \mid e_1 \times e_2 \mid f(e_1, \dots, e_k) \mid \text{for } v, X. e$$

Similarly as for `MATLANG`, we write `for-MATLANG` $[\mathcal{F}]$ for some set $\mathcal{F}$ of functions when these are required for the task at hand.

As a final example, we show that we can compute whether a graph contains a 4-clique using `for-MATLANG`.

Example 3.3. *To test for 4-cliques it suffices to consider the following expression with for-loops nested four times:*

$$\begin{aligned} &\text{for } u, X_1. X_1 + \\ &\quad \text{for } v, X_2. X_2 + \end{aligned}$$

$$\begin{aligned}
& \text{for } w, X_3. X_3 + \\
& \quad \text{for } x, X_4. X_4 + \\
& \quad \quad u^T \cdot V \cdot v \cdot u^T \cdot V \cdot w \cdot u^T \cdot V \cdot x \cdot \\
& \quad \quad v^T \cdot V \cdot w \cdot v^T \cdot V \cdot x \cdot w^T \cdot V \cdot x \cdot g(u, v, w, x)
\end{aligned}$$

with $g(u, v, w, x) = f(u, v) \cdot f(u, w) \cdot f(u, x) \cdot f(v, w) \cdot f(v, x) \cdot f(w, x)$ and $f(u, v) = 1 - u^T \cdot v$. Note that $f(b_i^n, b_j^n) = 1$ if $i \neq j$ and $f(b_i^n, b_j^n) = 0$ otherwise. Hence, $g(b_i^n, b_j^n, b_k^n, b_\ell^n) = 1$ if and only if all i, j, k, ℓ are pairwise different. When evaluating the expression on an instance $\mathcal{I}$ such that V is assigned to the adjacency matrix of a graph, the expression above evaluates to a non-zero value if and only if the graph contains a four-clique. $\square$

Given that `MATLANG` can not check for 4-cliques (Brijder et al., 2019), we easily obtain the following.

PROPOSITION 3.1. *For any collection of functions $\mathcal{F}$, $\text{MATLANG}[\mathcal{F}]$ is properly subsumed by $\text{for-MATLANG}[\mathcal{F}]$.*

3.2. Design decisions behind **for-MATLANG**

3.2.1. Loop Initialization

As the reader may have observed, in the semantics of for-loops we always initialize A_0 to the zero matrix $\mathbf{0}$ (of appropriate dimensions). It is often convenient to start the iteration given some concrete matrix originating from the result of evaluation a `for-MATLANG` expression e_0 . To make this explicit, we write `for` $v, X = e_0. e$ and its semantics is defined as above with the difference that $A_0 := \llbracket e_0 \rrbracket(\mathcal{I})$. We observe, however, that `for` $v, X = e_0. e$ can already be expressed in `for-MATLANG`. In other words, we do not loose generality by assuming an initialization of A_0 by $\mathbf{0}$. The key insight is that in `for-MATLANG` we can check during evaluation whether or not the current canonical vector b_i^n is equal to the b_1^n . This is due to the fact that for-loops iterate over the canonical

vectors in a fixed order. We discuss this more in the next paragraph. In particular, we can define a `for-MATLANG` expression `min()`, which when evaluated on an instance, returns 1 if its input vector is b_1^n , and returns 0 otherwise. Given `min()`, consider now the `for-MATLANG` expression

$$\text{for } v, X. \text{min}(v) \cdot e(v, X/e_0) + (1 - \text{min}(v)) \cdot e(v, X),$$

where we explicitly list v and X as matrix variables on which e potentially depends on, and where $e(v, X/e_0)$ denotes the expression obtained by replacing every occurrence of X in e with e_0 . When evaluating this expression on an instance $\mathcal{I}$, A_0 is initial set to the zero matrix, in the first iteration (when $v = b_1^n$ and thus $\text{min}(v) = 1$) we have $A_1 = \llbracket e \rrbracket(\mathcal{I}[v := b_1^n, X := \llbracket e_0 \rrbracket(\mathcal{I})])$, and for consecutive iterations (when only the part related to $1 - \text{min}(v)$ applies) A_i is updated as before. Clearly, the result of this evaluation is equal to $\llbracket \text{for } v, X = e_0. e \rrbracket(\mathcal{I})$.

As an illustration, we consider the Floyd-Warshall algorithm given in the Introduction.

Example 3.4. *Consider the following expression:*

$$\begin{aligned} e_{FW} := & \text{for } v_k, X_1 = A. X_1 + \\ & \text{for } v_i, X_2. X_2 + \\ & \text{for } v_j, X_3. X_3 + \\ & (v_i^T \cdot X_1 \cdot v_k \cdot v_k^T \cdot X_1 \cdot v_j) \times v_i \cdot v_j^T \end{aligned}$$

The expression e_{FW} simulates the Floyd-Warshall algorithm by updating the matrix A , which is stored in the variable X_1 . The inner sub-expression here constructs an $n \times n$ matrix that contains one in the position (i, j) if and only if one can reach the vertex j from i by going through k , and zero elsewhere. If an instance $\mathcal{I}$ assigns to A the adjacency matrix of a graph, then $\llbracket e_{FW} \rrbracket(\mathcal{I})$ will be equal to the matrix produced by the algorithm given in the Introduction. $\square$

3.2.2. Order

By introducing for-loops we not only extend MATLANG with bounded recursion, we also introduce order information. Indeed, the semantics of the `for` operator assumes that the canonical vectors $b_1, b_2, \dots$ are accessed in this order. It implies, among other things, that `for`-MATLANG expressions are not permutation-invariant. We can, for example, return the bottom right-most entry in a matrix. Indeed, consider the expression $e_{\max} := \text{for } v, X . v$ which, for it to be well-typed, requires both v and X to be of type $(\alpha, 1)$. Then, $\llbracket e_{\max} \rrbracket(\mathcal{I}) = b_n^n$, for $n = \mathcal{D}(\alpha)$, simply because initially, $X = \mathbf{0}$, but X will be overwritten by $b_1^n, b_2^n, \dots, b_n^n$, in this order. Hence, at the end of the evaluation b_n^n is returned. To extract the bottom right-most entry from a matrix, we now simply use $e_{\max}^T \cdot V \cdot e_{\max}$.

Although the order is implicit in `for`-MATLANG, we can explicitly use this order in `for`-MATLANG expressions. More precisely, the order on canonical vectors is made accessible by using the matrix:

$$S_{\leq} = \begin{bmatrix} 1 & 1 & \cdots & 1 \\ 0 & 1 & \cdots & 1 \\ \vdots & \vdots & \ddots & 1 \\ 0 & 0 & \cdots & 1 \end{bmatrix}.$$

We observe that $S_{\leq}$ has the property that $b_i^T \cdot S_{\leq} \cdot b_j = 1$, for two canonical vectors b_i and b_j of the same dimension, if and only if $i \leq j$. Otherwise, $b_i^T \cdot S_{\leq} \cdot b_j = 0$. Interestingly, we can build the matrix $S_{\leq}$ with the following `for`-MATLANG expression:

$$e_{\leq} = \text{for } v, X . X + ((X \cdot e_{\max}) + v) \cdot v^T + v \cdot e_{\max}^T,$$

where $e_{\max}$ is as defined above. The intuition behind this expression is that by using the last canonical vector b_n , as returned by $e_{\max}$, we have access to the last column of X (via the product $X \cdot e_{\max}$). We use this column such that after the i -th iteration, this column contains the i -th column of $S_{\leq}$. This is done by incrementing X with $v \cdot e_{\max}^T$. To construct

$S_{\leq}$, in the i -th iteration we further increment X with (i) the current last column in X (via $X \cdot e_{\max} \cdot v^T$) which holds the $(i-1)$ -th column of $S_{\leq}$; and (ii) the current canonical vector (via $v \cdot v^T$). Hence, after iteration i , X contains the first i columns of $S_{\leq}$ and holds the i th column of $S_{\leq}$ in its last column. It is now readily verified that $X = S_{\leq}$ after the n th iteration.

It should be clear that if we can compute $S_{\leq}$ using $e_{\leq}$, then we can easily define the following predicates and vectors related with the order of canonical vectors:

- $\text{succ}(u, v)$ such that $\text{succ}(b_i^n, b_j^n) = 1$ if $i \leq j$ and 0 otherwise. Similarly, we can define $\text{succ}^+(u, v)$ such that $\text{succ}^+(b_i^n, b_j^n) = 1$ if $i < j$ and 0 otherwise;
- $\text{min}(u)$ such that $\text{min}(b_i^n) = 1$ if $i = 1$ and $\text{min}(b_i^n) = 0$ otherwise;
- $\text{max}(u)$ such that $\text{max}(b_i^n) = 1$ if $i = n$ and $\text{min}(b_i^n) = 0$ otherwise; and
- $e_{\min}$ and $e_{\max}$ such that $\llbracket e_{\min} \rrbracket(\mathcal{I}) = b_1^n$ and $\llbracket e_{\max} \rrbracket(\mathcal{I}) = b_n^n$, respectively.

We can define these and many other order related expressions as follows.

By defining

$$\text{succ}(u, v) := u^T \cdot S_{\leq} \cdot v,$$

we obtain an order relation that allows us to discern whether one canonical vector comes before the other in the order given by $S_{\leq}$. If we want a strict order, we can just use the matrix $S_{<} := S_{\leq} - e_{\text{Id}}$, where e_{Id} is an expression in `for-MATLANG` which returns the identity matrix (of appropriate dimension). Given this, we define

$$\text{succ}^+(u, v) := u^T \cdot S_{<} \cdot v.$$

from which we can also derive

$$\text{max}(u) := u^T \cdot e_{\max}.$$

which is an expression that returns the last canonical vector.

Interestingly, we can also define the *previous* relation between canonical vectors. For this, we require the following matrix:

$$\text{Prev} = \begin{bmatrix} 0 & 1 & \cdots & 0 \\ 0 & \ddots & \ddots & \vdots \\ \cdots & \cdots & \cdots & 1 \\ 0 & \cdots & \cdots & 0 \end{bmatrix},$$

Using this matrix, we have that for a canonical vector b_i :

$$\text{Prev} \cdot b_i = \begin{cases} b_{i-1}, & \text{if } i > 1. \\ \mathbf{0}, & \text{if } i = 1. \end{cases}$$

where $\mathbf{0}$ is a vector of zeros of the same type as b_i . Notice also that $\mathbf{1}(u)^T \cdot \text{Prev} \cdot u$ is equal to zero, for a canonical vector u , if and only if $u = b_1$ is the first canonical vector, and zero otherwise. Therefore the expression $\min(u)$ is defined as

$$\min(u) := 1 - \mathbf{1}(u)^T \cdot \text{Prev} \cdot u,$$

and, when evaluated over canonical vectors, will result in 1 if and only if $u = b_1$ is the first canonical vector. To define the first canonical vector in the order given by `for`, we can then write:

$$e_{\min} := \text{for } v, X \cdot X + \min(v) \times v,$$

Now, we show that `Prev` can be defined using the following `for-MATLANG` expression:

$$e_{\text{Prev}} := \text{for } v, X \cdot X + ((1 - \max(v)) \times v \cdot e_{\max}^T - (X \cdot e_{\max}) \cdot e_{\max}^T + (X \cdot e_{\max}) \cdot v^T).$$

Here, X is initialized as $\mathbf{0}$ and thus in the first iteration we put b_1 in the last column of X (note that $X \cdot e_{\max}$ is also zero in the first iteration). Next, in iteration two, we add a matrix that has the stored vector $X \cdot e_{\max}$ (the previous canonical vector) in the column indicated by v (the current canonical vector) and $v - X \cdot e_{\max}$ in the last column, to replace the vector stored. As a consequence, b_2 is now stored in the last column. In the last iteration, we have b_{n-1} already in the last column, so no further update of X is required.

To get the *next* relation we simply do $e_{\text{Next}} = e_{\text{Prev}}^T$. We have that for a canonical vector b_i :

$$\text{Next} \cdot b_i = \begin{cases} b_{i+1}, & \text{if } i < n. \\ \mathbf{0}, & \text{if } i = n. \end{cases}$$

In this way, we also can obtain the following operators for a canonical vector v :

$$\text{prev}(v) := e_{\text{Prev}} \cdot v.$$

$$\text{next}(v) := e_{\text{Next}} \cdot v.$$

More generally, we define

$$e_{\text{getPrevMatrix}}(v) := \Pi w. \text{succ}(w, v) \times e_{\text{Prev}} + (1 - \text{succ}(w, v)) \times e_{\text{Id}}$$

$$e_{\text{getNextMatrix}}(v) := \Pi w. \text{succ}(w, v) \times e_{\text{Next}} + (1 - \text{succ}(w, v)) \times e_{\text{Id}}$$

expressions that, when v is interpreted as canonical vector b_i , output Prev^i and Next^i respectively. Note that

$$\text{Prev}^j \cdot b_i = \begin{cases} b_{i-j}, & \text{if } i > j. \\ \mathbf{0}, & \text{if } i \leq j. \end{cases}$$

and

$$\text{Next}^j \cdot b_i = \begin{cases} b_{i+j}, & \text{if } i + j \leq n. \\ \mathbf{0}, & \text{if } i + j > n. \end{cases}$$

Finally, define

$$e_{\min+i} := \underbrace{e_{\text{getNextMatrix}}(\dots e_{\text{getNextMatrix}}(e_{\min}))}_{i \text{ times}}$$

and

$$e_{\max-i} := \underbrace{e_{\text{getPrevMatrix}}(\dots e_{\text{getPrevMatrix}}(e_{\max}))}_{i \text{ times}}.$$

Having order information available results in `for-MATLANG` to be quite expressive. We heavily rely on order information in the next chapters to compute the inverse of matrices and more generally to simulate low complexity Turing machines and arithmetic circuits.

3.3. Syntactic definitions and simplifications

Before jumping into the following chapters we first introduce some additional notations and describe simplifications that will be used later.

3.3.1. Operators

Notice that when defining some expressions, we actually only update X by adding some matrix to it. This restricted form of the `for` loop will prove useful throughout the paper, and we will therefore introduce it a special operator. That is, we define:

$$\Sigma v.e := \text{for } v, X . X + e.$$

Analogously to the summation with respect to canonical vectors, we can define the product. More precisely, if e is an `for-MATLANG` expression (that possibly uses the variable v), we would like to define the result of evaluating $\llbracket e \rrbracket(\mathcal{I}[v := e_1]) \cdot \llbracket e \rrbracket(\mathcal{I}[v := e_2]) \cdots \llbracket e \rrbracket(\mathcal{I}[v := e_n])$, where the product is evaluated from left to right. For this purpose, we define the following operator:

$$\Pi v.e = \text{for } v, X . X \cdot e + \min(v) \cdot e.$$

Here we use the factor $\min(v) \cdot e$ to handle the case of the first canonical vector which starts with X being equal to the null matrix.

We sometimes want to iterate over k canonical vectors. We define the following shorthand notation:

$$\begin{aligned} \text{for } v_1, \dots, v_k, X . e(X, v_1, \dots, v_n) := & \text{for } v_1, X_1 . X_1 + \\ & \text{for } v_2, X_2 = X_1 . X_2 + \\ & \text{for } v_3, X_3 = X_2 . X_3 + \\ & \vdots \\ & \text{for } v_k, X_k = X_{k-1} . e(X_k, v_1, \dots, v_k). \end{aligned}$$

Here, we just aggregate the results of $e(X_k, v_1, \dots, v_k)$ for it to be returned after $v_1, \dots, v_k$ iterate over the canonical vectors. Note that the typing enforces $e, X_1, \dots, X_k$ to be of the same type.

To reference ℓ different vector variables $X_1, \dots, X_\ell$ in every iteration and update them in different ways we define:

$$\begin{aligned} \text{for } v, X_1, \dots, X_\ell. (e_1(X_1, v), e_2(X_2, v), \dots, e_\ell(X_\ell, v)) := \\ \text{for } v, X. e_1(X \cdot e_{\min}, v) \cdot (e_{\text{diag}}(e_1(X^T)) \cdot e_{\min})^T + \\ e_2(X \cdot e_{\min+1}, v) \cdot (e_{\text{diag}}(e_1(X^T)) \cdot e_{\min+1})^T + \dots + \\ e_\ell(X \cdot e_{\max}, v) \cdot (e_{\text{diag}}(e_1(X^T)) \cdot e_{\max})^T \end{aligned}$$

We note that for the latter expression to be semantically correct v has to be of type $\gamma \times 1$, both X_i and e_i for $i = 1, \dots, \ell$ have to be of type $\alpha \times 1$, and X has to be of type $\alpha \times \beta$, where $\mathcal{D}(\beta) = \ell$. Here we use $e_{\text{diag}}(e_1(X^T))$ to compute the $\beta \times \beta$ identity and ensure the typing of the $e_{\min+i}$. When evaluated on an instance $\mathcal{I}$, $e_{\min}, e_{\min+i}$ evaluate to $b_1^{\mathcal{D}(\beta)}$ and $b_{1+i}^{\mathcal{D}(\beta)}$, respectively, their defining expressions where shown in Chapter 3.2.2. Similarly for $e_{\max} = b_n^{\mathcal{D}(\beta)}$. The combinations of both previous operators results in:

$$\begin{aligned} \text{for } v_1, \dots, v_k, X_1, \dots, X_\ell. \\ e_1(X_1, v_1, \dots, v_k), e_2(X_2, v_1, \dots, v_k), \dots, e_\ell(X_\ell, v_1, \dots, v_k) := \\ \text{for } v_1, \dots, v_k, X. e'(X, v_1, \dots, v_k) \end{aligned}$$

where

$$\begin{aligned} e'(X, v_1, \dots, v_k) := & e_1(X \cdot e_{\min}, v_1, \dots, v_k) \cdot (e_{\text{diag}}(e_1(X^T)) \cdot e_{\min})^T \\ & + e_2(X \cdot e_{\min+1}, v_1, \dots, v_k) \cdot (e_{\text{diag}}(e_1(X^T)) \cdot e_{\min+1})^T \\ & + \dots + e_\ell(X \cdot e_{\max}, v_1, \dots, v_k) \cdot (e_{\text{diag}}(e_1(X^T)) \cdot e_{\max})^T \end{aligned}$$

It is clear that this expression iterates over k canonical vectors and references ℓ independent vectors updating each of them in their particular way.

3.3.2. Function application simplifications

When showing results based on induction of expressions in `for-MATLANG`, it is often convenient to assume that function applications $f(e_1, \dots, e_k)$ for $f \in \mathcal{F}_k$ are restricted to the case when all expressions $e_1, \dots, e_k$ have type 1×1 . This does not loose generality. Indeed, for general function applications $f(e_1, \dots, e_k)$, if we have Σ (introduced in Chapter 6.1), scalar product and function application on scalars (here denoted by $f_{1 \times 1}$), we can simulate full function application, as follows:

$$f(e_1, \dots, e_k) := \Sigma v_i \Sigma v_j. f_{1 \times 1}(v_i^T \cdot e_1 \cdot v_j, \dots, v_i^T \cdot e_k \cdot v_j) \times v_i \cdot v_j^T.$$

Furthermore, it is also convenient at times to use the pointwise functions $f_{\odot}^k : \mathbb{R}^k \mapsto \mathbb{R} : (x_1, \dots, x_k) \mapsto x_1 \times \dots \times x_k$ and $f_{\oplus}^k : \mathbb{R}^k \mapsto \mathbb{R} : (x_1, \dots, x_k) \mapsto x_1 + \dots + x_k$. In fact, it is readily observed that adding these functions does not extend the expressive power of `for-MATLANG`:

Lemma 3.1. *We have that $\text{for-MATLANG}[\emptyset] \equiv \text{for-MATLANG}[\{f_{\odot}^k, f_{\oplus}^k \mid k \in \mathbb{N}\}]$.*

In fact, this lemma also holds for the smaller fragments we consider.

We also observe that having $f_{\odot}^2 : \mathbb{R}^2 \rightarrow \mathbb{R}$ allows us to define scalar multiplication:

$$e_1 \times e_2 := f_{\odot}(\mathbf{1}(e_2)^T \cdot e_1 \cdot \mathbf{1}(e_2)^T, e_2).$$

Conversely, $f_{\odot}^k$ can be expressed using scalar multiplication, as can be seen from our simulation of general function applications by pointwise function application on scalars.

Finally, a notational simplification is that when using scalars $a \in \mathbb{R}$ in our expressions, we write sometimes a instead of $[a]$, because everything is a matrix in this context. For example, $(1 - e_1(v)^T \cdot v)$ stands for $([1] - e_1(v)^T \cdot v)$.

4. ALGORITHMS IN LINEAR ALGEBRA

One of our main motivations to introduce for-loops is to be able to express classical linear algebra algorithms in a natural way. We have seen that `for-MATLANG` is quite expressive as it can check for cliques, compute the transitive closure, and can even leverage a successor relation on canonical vectors. The big question is how expressive `for-MATLANG` actually is. We will answer this in the next chapter by connecting `for-MATLANG` with arithmetic circuits of polynomial degree. Through this connection, one can move back and forth between `for-MATLANG` and arithmetic circuits, and as a consequence, anything computable by such a circuit can be computed by `for-MATLANG` as well. When it comes to specific linear algebra algorithms, the detour via circuits can often be avoided. Indeed, in this chapter we illustrate that `for-MATLANG` is able to compute LU decompositions of matrices. These decompositions form the basis of many other algorithms, such as solving linear systems of equations. We further show that `for-MATLANG` is expressive enough to compute matrix inversion and the determinant. We recall that matrix inversion and determinant need to be explicitly added as separate operators in `MATLANG` (Brijder et al., 2019) and that the LARA language is unable to invert matrices under usual complexity-theoretic assumptions (Barceló et al., 2020).

4.1. LU decomposition

A lower-upper (LU) decomposition factors a matrix A as the product of a lower triangular matrix L and upper triangular matrix U . This decomposition, and more generally LU decomposition with row pivoting (PLU), underlies many linear algebra algorithms and we next show that `for-MATLANG` can compute these decompositions.

4.1.1. LU decomposition by Gaussian elimination

LU decomposition can be seen as a matrix form of Gaussian elimination in which the columns of A are reduced, one by one, to obtain the matrix U . The reduction of

columns of A is achieved as follows. Consider the first column $[A_{11}, \dots, A_{n1}]^T$ of A and define $c_1 := [0, \alpha_{21}, \dots, \alpha_{n1}]^T$ with $\alpha_{j1} := -\frac{A_{j1}}{A_{11}}$. Let $T_1 := I + c_1 \cdot b_1^T$ and consider $T_1 \cdot A$. That is, the j th row of $T_1 \cdot A$ is obtained by multiplying the first row of A by α_{j1} and adding it to the j th row of A . As a result, the first column of $T_1 \cdot A$ is equal to $[A_{11}, 0, \dots, 0]^T$, i.e., all of its entries below the diagonal are zero. One then iteratively performs a similar computation, using a matrix $T_i := I + c_i \cdot b_i^T$, where c_i now depends on the i th column in $T_{i-1} \cdots T_1 \cdot A$. As a consequence, $T_i \cdot T_{i-1} \cdots T_1 \cdot A$ is upper triangular in its first i columns. At the end of this process, $T_n \cdots T_1 \cdot A = U$ where U is the desired upper triangular matrix. Furthermore, it is easily verified that each T_i is invertible and by defining $L := T_1^{-1} \cdots T_n^{-1}$ one obtains a lower triangular matrix satisfying $A = L \cdot U$. The above procedure is only successful when the denominators used in the definition of the vectors c_i are non-zero. When this is the case we call a matrix A *LU-factorizable*.

In case when such a denominator is zero in one of the reduction steps, one can remedy this situation by *row pivoting*. That is, when the i th entry of the i th row in $T_{i-1} \cdots T_1 \cdot A$ is zero, one replaces the i th row by j th row in this matrix, with $j > i$, provided that the entry of the j th row is non-zero. If no such row exists, this implies that all elements below the diagonal are zero already in column i and one can proceed with the next column. One can formulate this in matrix terms by stating that there exists a permutation matrix P , which pivots rows, such that $P \cdot A = L \cdot U$. Any matrix A is LU-factorizable *with pivoting*.

4.1.2. Implementing LU decomposition in **for-MATLANG**

We first assume that the input matrices are LU-factorizable. We deal with general matrices later on. To implement the above procedure, we need to compute the vector c_i for each column i . We do this in two steps. First, we extract from our input matrix its i th column and set all its upper diagonal entries to zero by means of the expression:

$$\text{col}(V, y) := \text{for } v, X. \text{succ}^+(y, v) \cdot (v^T \cdot V \cdot y) \cdot v + X.$$

Indeed, when V is assigned to a matrix A and y to b_i , we have that X will be initially assigned $A_0 = \mathbf{0}$ and in consecutive iterations, $A_j = A_{j-1} + b_j^T \cdot A \cdot b_i$ if $j > i$ (because $\text{succ}^+(b_i, b_j) = 1$ if $j > i$) and $A_j = A_{j-1}$ otherwise (because $\text{succ}^+(b_i, b_j) = 0$ for $j \leq i$). The result of this evaluation is the desired column vector. Using $\text{col}(V, y)$, we can now compute T_i by the following expression:

$$\text{reduce}(V, y) := e_{\text{ld}} + f_{/}(\text{col}(V, y), -(y^T \cdot V \cdot y) \cdot \mathbf{1}(y)) \cdot y^T,$$

where $f_{/} : \mathbb{R}^2 \rightarrow \mathbb{R} : (x, y) \mapsto x/y$ is the division function. When V is assigned to A and y to b_i , $f_{/}(\text{col}(A, b_i), -(b_i^T \cdot A \cdot b_i) \cdot \mathbf{1}(b_i))$ is equal to the vector c_i used in the definition of T_i . To perform the reduction steps for all columns, we consider the expression:

$$e_U(V) := (\text{for } y, X = e_{\text{ld}} \cdot \text{reduce}(X \cdot V, y) \cdot X) \cdot V.$$

That is, when V is assigned A , X will be initially $A_0 = I$, and then $A_i = \text{reduce}(A_{i-1} \cdot A, b_i) = T_i \cdot T_{i-1} \cdots T_1 \cdot A$, as desired. We show that, because we can obtain the matrices T_i in `for-MATLANG` and that these are easily invertible, we can also construct an expression $e_L(V)$ which evaluates to L when V is assigned to A . We may thus conclude the following.

PROPOSITION 4.1. *There exists `for-MATLANG`[$f_{/}$] expressions $e_L(V)$ and $e_U(V)$ such that $\llbracket e_L \rrbracket(\mathcal{I}) = L$ and $\llbracket e_U \rrbracket(\mathcal{I}) = U$ form an LU-decomposition of A , where $\text{mat}(V) = A$ and A is LU-factorizable.*

PROOF. Let A be an LU-factorizable matrix. We already explained how the expression $e_U(V)$ is obtained, i.e.,

$$e_U(V) := (\text{for } y, X = e_{\text{ld}} \cdot \text{reduce}(X \cdot V, y) \cdot X) \cdot V.$$

We recall that $e_U(A) = T_n \cdots T_1 \cdot A$ with $L^{-1} = T_n \cdots T_1$. Let

$$e_{L^{-1}}(V) := \text{for } y, X = e_{\text{ld}} \cdot \text{reduce}(X \cdot V, y) \cdot X.$$

such that

$$e_U(V) := e_{L^{-1}}(V) \cdot V.$$

It now suffices to observe that, since $T_n = I$,

$$\begin{aligned} L^{-1} &= (I - c_1 \cdot b_1^T) \cdots (I - c_{n-1} \cdot b_{n-1}^T) \\ &= I - c_1 \cdot b_1^T - \cdots - c_{n-1} \cdot b_{n-1}^T \end{aligned}$$

and hence,

$$\begin{aligned} L &= (I + c_1 \cdot b_1^T) \cdots (I + c_{n-1} \cdot b_{n-1}^T) \\ &= I + c_1 \cdot b_1^T + \cdots + c_{n-1} \cdot b_{n-1}^T. \end{aligned}$$

As a consequence, to obtain L from L^{-1} we just need to multiply every entry below the diagonal by -1 . Since both L and L^{-1} are lower triangular, this can be done by computing $L = -1 \times L^{-1} + 2 \times I$. Translated into `for-MATLANG`, this means that we can define

$$e_L(V) := -1 \times e_{L^{-1}}(V) + 2 \times e_{\text{Id}},$$

which concludes the proof of the proposition. $\square$

We remark that the proposition holds when division is added as a function in $\mathcal{F}$ in `for-MATLANG`. When row pivoting is needed, we can also obtain a permutation matrix P such that $P \cdot A = L \cdot U$ holds by means of an expression in `for-MATLANG`, provided that we additionally allow the function $f_{>0}$, where $f_{>0} : \mathbb{R} \rightarrow \mathbb{R}$ is such that $f_{>0}(x) := 1$ if $x > 0$ and $f_{>0}(x) := 0$ otherwise. Intuitively, by allowing $f_{>0}$ we introduce a limited form of `if-then-else` in `for-MATLANG`, which is needed to continue reducing columns only when the right pivot has been found.

PROPOSITION 4.2. *There exist expressions $e_{L^{-1}P}(M)$ and $e_U(M)$ in `for-MATLANG` $[f_{\neq}, f_{>0}]$ such that $L^{-1} \cdot P = \llbracket e_{L^{-1}P} \rrbracket(\mathcal{I})$ and $U = \llbracket e_U \rrbracket(\mathcal{I})$, satisfy $L^{-1} \cdot P \cdot A = U$.*

PROOF. We assume that $f_{\neq}$ and $f_{>0}$ are in $\mathcal{F}$. Let A be an arbitrary matrix. By contrast to when A is LU-factorizable, during the LU-decomposition process we may need row interchange (pivoting) in each step of the iteration. Let us assume that row interchange is needed immediately before step k , $1 \leq k \leq n$. In other words, we now aim to reduce the k -th column of $A_k = T_{k-1} \cdots T_1 \cdot A$, or $A_k = A$ if $k = 1$, but now A_k has a zero pivot,

i.e., $(A_k)_{kk} = 0$. Let P be the matrix that denotes the necessary row interchange. If we know P , then to compute T_k we need to perform $\text{reduce}(P \cdot X \cdot A, v)$ in this iteration, where $\text{reduce}(\cdot, \cdot)$ is the expression in `for-MATLANG` reducing a column, as explained earlier. Furthermore, we need to apply the permutation P to the current result, resulting in the expression `for  $v, X = I \cdot \text{reduce}(P \cdot X \cdot A, v) \cdot P \cdot X$` . We now remark that P is a permutation matrix of the form $P = I - u \cdot u^T$ and it denotes an interchange (if multiplied by left) of rows i and j if $u = (b_i - b_j)$. Note that we are performing a row interchange for column k and thus $i = k$ and $j > k - 1$. If no interchange is needed, $i = j = k$ and $P = I$. Also note that when $k = n$ no interchange takes place. Furthermore, if no suitable b_j can be found, this implies that no interchange is required as well and we can move on to next column.

To find the vector u in P , we can, for example, find the first entry $j \geq k$ in column k of A_k that holds a non-zero value. More generally, we can find the first entry in a vector a that holds a non-zero value by using the function $f_{>0}$. Indeed, consider the following expression:

$$\begin{aligned} \text{neq}(a, u) := & \text{for } v, X \cdot (1 - e_1(v)^T \cdot X) \times f_{>0}((v^T \cdot a)^2) \times v + \\ & \max(v) \times (1 - e_1(v)^T \cdot X) \times (1 - f_{>0}((v^T \cdot a)^2)) \times u \end{aligned}$$

Here, $\text{neq}(a, u)$ receives two n dimensional vectors a and u and outputs a canonical vector b_j such that a_j is the first non-zero entry of a , or u if such non-zero value does not exist in a . We check for $f_{>0}((\cdot)^2)$ in case a negative number is tested. The above expression simply checks in each iteration whether X already holds a canonical vector. If so, then X is not updated. Otherwise, X is replaced by the current canonical vector b_j if and only if $b_j^T \cdot a$ is non-zero. Furthermore, when the final canonical vector is considered and X does not hold a canonical vector yet and $b_n^T \cdot a$ is zero, the vector u is returned.

We use $\text{neq}(a, u)$ to find a pivot for a specific column. Let us assume again that we want to find a pivot in column k of A_k . We can then first make all entries in that column, with indexes smaller or equal to k , zero, just as we did by means of `col( $\cdot, \cdot$ )` in the

definition of $\text{reduce}(\cdot, \cdot)$. Except, now we also need to make the k th entry zero as well. Let us denote by $\text{col}_{\text{eq}}(\cdot, \cdot)$ the operation $\text{col}(\cdot, \cdot)$, as defined before, but using succ instead of succ^+ (to include the k entry). Given this, we can construct $P = I - u \cdot u^T$ as follows:

$$e_{P_u}(A, u) := e_{\text{Id}} - [u - \text{neq}(\text{col}_{\text{eq}}(A, u), u)] \cdot [u - \text{neq}(\text{col}_{\text{eq}}(A, u), u)]^T.$$

From the explanations given above, it should be clear that $e_{P_u}(A, u)$ computes the necessary permutation matrix of A_k for the column indicated by u , or I if no permutation is needed, or if such permutation does not exist (so we skip the current column). Also, we have to modify the $\text{reduce}(V, y)$ operators, as follows:

$$\begin{aligned} \text{reduce}(V, y) := & e_{\text{Id}} + f_{>0}((y^T \cdot V \cdot y)^2) \times f_l(\text{col}(V, y), \\ & [-(y^T \cdot V \cdot y) \times e_1(y) + (1 - f_{>0}((y^T \cdot V \cdot y)^2)) \times e_1(y)]) \cdot y^T, \end{aligned}$$

so that when V is interpreted by a matrix B and $y = b_i$, it returns $I + c_i b_i^T$ if B_{ii} is not zero. If $B_{ii} = 0$ then we divide $\text{col}(B, b_i)$ by $e_1(b_i)$ (so we don't get *undefined*), but we don't add $c_i b_i^T$ precisely because $B_{ii} = 0$, and return the identity so nothing happens. We check for $f_{>0}((\cdot)^2)$ in case a negative number is tested.

Finally, we define

$$e_{L^{-1}P}(V) := \text{for } v, X = e_{\text{Id}} \cdot \text{reduce}(e_{P_v}(X \cdot V, v) \cdot X \cdot V, v) \cdot e_{P_v}(X \cdot V, v) \cdot X$$

and $e_U(V) := e_{L^{-1}P}(V) \cdot V$ as the desired expressions.

As a final observation, in the definition of $e_{L^{-1}P}(V)$ we interlaced permutation matrices with the T_i 's. More specifically, $A_k = T_k \cdot P \cdot T_{k-1} \cdots T_1 \cdot A$. We observe, however, that for $\ell \leq k-1$ and $T_\ell = I - c_\ell \cdot b_\ell^T$, we have that $b_\ell^T \cdot P = b_\ell^T$ because b_ℓ has zeroes in positions in the rows involved in the row exchange P . Also, note that $P^2 = I$ and thus

$$P \cdot T_\ell \cdot P = P^2 - P \cdot c_\ell \cdot b_\ell^T \cdot P = I - \widehat{c}_\ell \cdot b_\ell^T = \widehat{T}_\ell.$$

Where $\widehat{c}_\ell = P \cdot c_\ell$.

As a consequence,

$$\begin{aligned}
T_k \cdot P \cdot T_{k-1} \cdots T_1 &= T_k \cdot P \cdot T_{k-1} \cdot P^2 \cdot T_{k-2} \cdot P^2 \cdots P^2 \cdot T_1 \cdot P^2 \\
&= T_k \cdot (P \cdot T_{k-1} \cdot P) \cdots (P \cdot T_1 \cdot P) \cdot P \\
&= \hat{T}_{k-1} \cdots \hat{T}_1 \cdot P,
\end{aligned}$$

and thus we may assume that P occurs at the end. Hence, we obtain $L^{-1} \cdot P \cdot A = U$. $\square$

4.2. Determinant and inverse

Other key linear algebra operations include the computation of the determinant and the inverse of a matrix (if the matrix is invertible). As a consequence of the expressibility in $\text{for-MATLANG}[f/, f_{>0}]$ of LU-decompositions with pivoting, it can be shown that the determinant and inverse can be expressed as well.

However, the results in the next chapter (connecting for-MATLANG with arithmetic circuits) imply that the determinant and inverse of a matrix can already be defined in $\text{for-MATLANG}[f/]$. So instead of using LU decomposition with pivoting for matrix inversion and computing the determinant, we provide an alternative solution.

More specifically, we rely on Csanky's algorithm for computing the inverse of a matrix (Csanky, 1976).

To show Proposition 4.3 we first show that it holds when considering non-singular lower or upper triangular matrices.

Lemma 4.1. *There are $\text{for-MATLANG}[f/]$ expressions $e_{\text{upperDiagInv}}(V)$ and $e_{\text{lowerDiagInv}}(V)$ such that $\llbracket e_{\text{upperDiagInv}} \rrbracket(\mathcal{I}) = A^{-1}$ when $\mathcal{I}$ assigns V to an invertible upper triangular matrix A and $\llbracket e_{\text{lowerDiagInv}} \rrbracket(\mathcal{I}) = A^{-1}$ when $\mathcal{I}$ assigns V to an invertible lower triangular matrix A .*

PROOF. We start by considering the expression:

$$e_{\text{ps}}(V) := e_{\text{Id}} + \Sigma v. \Pi w. [\text{succ}(w, v) \times V + (1 - \text{succ}(w, v)) \times e_{\text{Id}}].$$

Here, $e_{\text{ps}}(A)$ results in $I + A + A^2 + \dots + A^n$ for any matrix A . In the expression, the outer loop defines which power we compute. That is, when v is the i th canonical vector, we compute A^i . Computing A^i is achieved via the inner product loop, which uses $\text{succ}(w, v)$ to determine whether w comes before or is v in the ordering of canonical vectors. When this is the case, we multiply the current result by A , and when w is greater than v , we use the identity as not to affect the already computed result. We add the identity at the end.

Now, let A be an $n \times n$ matrix that is upper triangular and let D_A be the matrix consisting of the diagonal elements of A , i.e.,

$$D_A = \begin{bmatrix} a_{11} & \cdots & \cdots & 0 \\ 0 & a_{22} & \cdots & 0 \\ 0 & \ddots & \vdots & \vdots \\ \vdots & \cdots & \cdots & a_{nn} \end{bmatrix}.$$

We can compute D_A by the expression:

$$e_{\text{getDiag}}(V) := \Sigma v. (v^T \cdot V \cdot v) \times v \cdot v^T.$$

Let $T = A - D_A$, then

$$A^{-1} = [D_A + T]^{-1} = [D_A (I + D_A^{-1}T)]^{-1} = (I + D_A^{-1}T)^{-1} D_A^{-1}.$$

We now observe that D_A^{-1} simply consists of the inverses of the elements on the diagonal. This can be expressed, as follows:

$$e_{\text{diagInverse}}(V) := \Sigma v. f_{/}(1, v^T \cdot e_{\text{getDiag}}(V) \cdot v) \times v \cdot v^T = \Sigma v. f_{/}(1, v^T \cdot V \cdot v) \times v \cdot v^T,$$

Where $f_{/}$ is the division function. In the last equality we take advantage of the fact that the diagonals of A and D_A are the same.

We now focus on the computation of $(I + D_A^{-1} \cdot T)^{-1}$. First, by construction, $D_A^{-1} \cdot T$ is strictly upper triangular and thus nilpotent, such that $(D_A^{-1} \cdot T)^n = 0$, where n is the dimension of A . Recall the following algebraic identity

$$(1 + x) \left(\sum_{i=0}^m (-x)^i \right) = 1 - (-x)^{m+1}.$$

By choosing $m = n - 1$ and applying it to $x = D_A^{-1} \cdot T$, we have

$$(I + D_A^{-1} \cdot T) \left(\sum_{i=0}^{n-1} (-D_A^{-1} \cdot T)^i \right) = I - (-D_A^{-1} \cdot T)^n = I.$$

Hence,

$$(I + D_A^{-1} \cdot T)^{-1} = \sum_{i=0}^{n-1} (-D_A^{-1} \cdot T)^i = \sum_{i=0}^n (-D_A^{-1} \cdot T)^i.$$

We now observe that

$$e_{\text{ps}}(-1 \times D_A^{-1} \cdot T) = \sum_{i=0}^n (-D_A^{-1} \cdot T)^i = (I + D_A^{-1} \cdot T)^{-1},$$

and thus

$$A^{-1} = e_{\text{ps}}(-1 \times [e_{\text{diagInverse}}(A)(A - e_{\text{getDiag}}(A))]) e_{\text{diagInverse}}(A).$$

Seeing this as an expression:

$$e_{\text{upperDiagInv}}(V) := e_{\text{ps}}(-1 \times [e_{\text{diagInverse}}(V)(V - e_{\text{getDiag}}(V))]) e_{\text{diagInverse}}(V),$$

we see that when interpreting V as an upper triangular invertible matrix, $e_{\text{upperDiagInv}}(A)$ evaluates to A^{-1} .

To deal with invertible lower triangular matrices A , we observe that $(A^{-1})^T = (A^T)^{-1}$ and A^T is upper triangular. Hence, it suffices to define

$$e_{\text{lowerDiagInv}}(V) := e_{\text{upperDiagInv}}(V^T)^T.$$

This concludes the proof of the lemma. □

We are now ready to conclude:

PROPOSITION 4.3. *There are $\text{for-MATLANG}[f_!]$ expressions $e_{\det}(V)$ and $e_{\text{inv}}(V)$ such that $\llbracket e_{\det} \rrbracket(\mathcal{I}) = \det(A)$, and $\llbracket e_{\text{inv}} \rrbracket(\mathcal{I}) = A^{-1}$ when $\mathcal{I}$ assigns V to A and A is invertible.*

PROOF. Let A be an $n \times n$ matrix. As mentioned before, we will implement Csanky's algorithm. Let $p_A(x) := \det(xI - A)$ denote characteristic polynomial of A . We write $p_A(x) = 1 + \sum_{i=1}^n c_i x^i$ and let $S_i := \frac{1}{i+1} \text{tr}(A^i)$ with $\text{tr}(\cdot)$ the trace operator which sums up the diagonal elements of a matrix. Then, the coefficients $c_1, \dots, c_n$ are known to satisfy¹

$$\underbrace{\begin{pmatrix} 1 & 0 & 0 & \cdots & 0 & 0 \\ S_1 & 1 & 0 & \cdots & 0 & 0 \\ S_2 & S_1 & 1 & \cdots & 0 & 0 \\ \vdots & \vdots & \vdots & \ddots & \vdots & 0 \\ S_{n-1} & S_{n-2} & S_{n-3} & \cdots & S_1 & 1 \end{pmatrix}}_S \cdot \underbrace{\begin{pmatrix} c_1 \\ c_2 \\ c_3 \\ \vdots \\ c_n \end{pmatrix}}_{\vec{c}} = \underbrace{\begin{pmatrix} S_1 \\ S_2 \\ S_3 \\ \vdots \\ S_n \end{pmatrix}}_{\vec{b}}$$

and furthermore, $c_n = (-1)^n \det(A)$ and if $c_n \neq 0$, then

$$A^{-1} = \frac{1}{c_n} \sum_{i=0}^{n-1} c_i A^{n-1-i},$$

with $c_0 = 1$. It is now easy to see that we can compute the S'_i 's in for-MATLANG . Indeed, for $i = 1, \dots, n$ we can consider

$$e_{\text{powTr}}(V, v) := \Sigma w. w^T \cdot (e_{\text{pow}}(V, v) \cdot V) \cdot w$$

with

$$e_{\text{pow}}(V, v) := \Pi w. (\text{succ}(w, v) \times V + (1 - \text{succ}(w, v)) \times e_{\text{Id}}).$$

We have that $e_{\text{pow}}(A, b_j) = A^j$ and thus $e_{\text{powTr}}(A, b_j) = \text{tr}(A^j)$. Define:

$$e_S(V, v) := f_!(1, 1 + \Sigma w. \text{succ}(w, v)) \times e_{\text{powTr}}(V, v).$$

¹We use a slightly different, but equivalent, system of equations used by Csanky, 1976.

Here $e_S(A, b_i) = S_i$. Note that $i+1$ is computed summing up to the dimension indicated by v , and adding 1. We can now easily construct the vector $\bar{b}$ used in the system of equations by means of the expression:

$$e_{\bar{b}}(V) := \Sigma w.e_S(V, w) \times w.$$

We next construct the matrix S . We need to be able to *shift* a vector a in k positions, i.e., such that $(a_1, \dots, a_n) \mapsto (0, \dots, a_1, \dots, a_{n-k})$. We use $e_{\text{getNextMatrix}}$ defined in Chapter 3.2.2, i.e., we define:

$$e_{\text{shift}}(a, v) := \Sigma w.(w^T \cdot a) \times (e_{\text{getNextMatrix}}(v) \cdot w)$$

performs the desired shift when u is assigned a vector a and v is b_k . The matrix S is now obtained as follows:

$$S(V) := e_{\text{Id}} + \Sigma v.e_{\text{shift}}(e_{\bar{b}}(V), v) \cdot v^T$$

We now observe that S is lower triangular with nonzero diagonal entries. So, Lemma 4.1 tells us that we can invert it, i.e., $e_{\text{lowerDiagInv}}(S) = S^{-1}$. As a consequence,

$$e_{\bar{c}}(V) := e_{\text{lowerDiagInv}}(S(V)) \cdot e_{\bar{b}}(V).$$

outputs $\bar{c}$ when V is interpreted as matrix A . Observe that we only use the division operator. We now have all coefficients of the characteristic polynomial of A .

We can now define

$$e_{\text{det}}(V) := \left(\left((\Pi w.(-1) \times e_1(V))^T \cdot e_{\text{max}} \right) \times e_{\bar{c}}(V) \right)^T \cdot e_{\text{max}},$$

an expression that, when V is interpreted as any matrix A , outputs $\det(A)$. Here, $(\Pi w.(-1) \times e_1(V))$ is the n dimensional vector with $(-1)^n$ in all of its entries. Since $c_n = (-1)^n \det(A)$, we extract $(-1)^n(-1)^n \det(A) = \det(A)$ with e_{max} .

For the inverse, we have that

$$A^{-1} = \frac{1}{c_n} \sum_{i=0}^{n-1} c_i A^{n-1-i} = \frac{1}{c_n} A^{n-1} + \sum_{i=1}^{n-1} \frac{c_i}{c_n} A^{n-1-i}.$$

We compute $\frac{1}{c_n} A^{n-1}$ as

$$f_/(1, e_{\bar{c}}(A)^T \cdot e_{\max}) \times e_{\text{pow}}(A, e_{\max})$$

and $\sum_{i=1}^{n-1} \frac{c_i}{c_n} A^{n-1-i}$ as

$$\Sigma v.f_/(e_{\bar{c}}(A)^T \cdot v, e_{\bar{c}}(A)^T \cdot e_{\max}) \times e_{\text{invPow}}(A, v),$$

where

$$e_{\text{invPow}}(V, v) :=$$

$$\Pi w.(1 - \max(w)) \times [(1 - \text{succ}(w, v)) \times V + \text{succ}(w, v) \times e_{\text{Id}}] + \max(w) \times e_{\text{Id}}.$$

Here, $e_{\text{invPow}}(A, b_i) = A^{n-1-i}$ and $e_{\text{invPow}}(A, b_n) = I$. Note that we always multiply by e_{Id} in the last step. To conclude, we define:

$$e_{\text{inv}}(V) := f_/(1, e_{\bar{c}}(V)^T \cdot e_{\max}) \times e_{\text{pow}}(V, e_{\max}) + \\ [\Sigma v.f_/(e_{\bar{c}}(V)^T \cdot v, e_{\bar{c}}(V)^T \cdot e_{\max}) \times e_{\text{invPow}}(V, v)],$$

an expression that, when V is interpreted as any invertible matrix A , computes A^{-1} . $\square$

As an observation, here we only use operators Σ and Π defined in Chapter 3.3.1. We also assume access to order.

5. EXPRESSIVENESS OF FOR LOOPS

In this chapter we explore the expressive power of `for-MATLANG`. Given that arithmetic circuits (Allender, 2004) capture most standard linear algebra algorithms (Raz, 2003; Shpilka & Yehudayoff, 2010), they seem as a natural candidate for comparison. Intuitively, an arithmetic circuit is similar to a boolean circuit (Arora & Barak, 2009), except that it has gates computing the sum and the product function, and processes elements of $\mathbb{R}$ instead of boolean values. To connect `for-MATLANG` to arithmetic circuits we need a notion of uniformity of such circuits. After all, a `for-MATLANG` expression can take matrices of arbitrary dimensions as input and we want to avoid having different circuits for each dimension. To handle inputs of different sizes, we thus consider a notion of uniform families of arithmetic circuits, defined via a Turing machine generating a description of the circuit for each input size n .

What we show in the remainder of this chapter is that any function f which operates on matrices, and is computed by a uniform family of arithmetic circuits of bounded fan-in, can also be computed by a `for-MATLANG` expression, and vice versa. In order to keep the notation light, we will focus on `for-MATLANG` schemas over “square matrices” where each variable has type (α, α) , $(\alpha, 1)$, $(1, \alpha)$, or $(1, 1)$, although all of our results hold without these restrictions as well. In what follows, we will write `for-MATLANG` to denote `for-MATLANG` $[\emptyset]$, i.e. the fragment of our language with no additional pointwise functions. We begin by defining circuits and then show how circuit families can be simulated by `for-MATLANG`.

5.1. Linear space functions

We start by showing a crucial ingredient for making the correspondence between `for-MATLANG` and arithmetic circuits. More specifically, we show that any polynomial time Turing machine, working within linear space and producing linear space output, can

be simulated in `FOR-MATLANG`. For this proof and chapter only, we will denote the canonical vectors as $\mathbf{e}_1, \dots, \mathbf{e}_n$, since b will be used to represent a value on a position of a tape.

We consider deterministic Turing Machines (TM) T consisting of ℓ read-only input tapes, denoted by $R_1, \dots, R_\ell$, a work tape, denoted by W , and a write-only output tape, denoted by O . The TM T has a set Q of m states, denoted by $q_0, \dots, q_m$. We assume that q_0 is the initial state and q_m is the accepting state. The input and tape alphabet are $\Sigma = \{0, 1\}$ and $\Gamma = \Sigma \cup \{\triangleright, \triangleleft\}$, respectively. The special symbol $\triangleright$ denotes the beginning of each of the tapes, the symbol $\triangleleft$ denotes the end of the ℓ input tapes. The transition function Δ is defined as usual, i.e., $\Delta : Q \times \Gamma^{\ell+2} \rightarrow Q \times \Gamma^2 \times \{\leftarrow, \sqcup, \rightarrow\}^{\ell+2}$ such that $\Delta(q, (a_1, \dots, a_\ell, b, c)) = (q', (b', c'), (d_1, \dots, d_{\ell+2}))$ with $d_i \in \{\leftarrow, \sqcup, \rightarrow\}$, means that when T is in state q and the $\ell + 2$ heads on the tapes read symbols $a_1, \dots, a_\ell, b, c$, respectively, then T transitions to state q' , writes b', c' on the work and output tapes, respectively, at the position to which the work and output tapes' heads points at, and finally moves the heads on the tapes according $d_1, \dots, d_{\ell+2}$. More specifically, $\leftarrow$ indicates a move to the left, $\rightarrow$ a move to the right, and finally, $\sqcup$ indicates that the head does not move.

We assume that Δ is defined such that it ensures that on none of the tapes, heads can move beyond the leftmost marker $\triangleright$. Furthermore, the tapes $R_1, \dots, R_\ell$ are treated as read-only and the heads on these tapes cannot move beyond the end markers $\triangleleft$. Similarly, Δ ensures that the output tape O is write only, i.e., its head cannot move to the left. We also assume that Δ does not change the occurrences of $\triangleright$ or writes $\triangleleft$ on the work and output tape.

A configuration of T is defined in the usual way. That is, a configuration of the input tapes is of the form $\triangleright w_1 q w_2 \triangleleft$ with $w_1, w_2 \in \Sigma^*$ and represents that the current tape content is $\triangleright w_1 w_2 \triangleleft$, T is in state q and the head is positioned on the first symbol of w_2 . Similarly, configurations of the work and output tape are represented by $\triangleright w_1 q w_2$. A configuration of T consists of configurations for all tapes. Given two configurations c_1 and c_2 , we say that c_1 yields c_2 if c_2 is the result of applying the transition function Δ of T based on the information in c_1 . As usual, we close this “yields” relation transitively.

Given ℓ input words $w_1, \dots, w_\ell \in \Sigma^*$, we assume that the initial configuration of T is given by $(q_0 \triangleright w_1 \triangleleft, q_0 \triangleright w_2 \triangleleft, \dots, q_0 \triangleright w_\ell \triangleleft, q_0 \triangleright, q_0 \triangleright)$ and an accepting configuration is assumed to be of the form $(\triangleright q_m w_1 \triangleleft, \triangleright q_m w_2 \triangleleft, \dots, \triangleright q_m w_\ell \triangleleft, \triangleright q_m, \triangleright q_m w)$ for some $w \in \Sigma^*$. We say that T computes the function $f : (\Sigma^*)^\ell \rightarrow \Sigma^*$ if for every $w_1, \dots, w_\ell \in \Sigma^*$, the initial configuration yields (transitively) an accepting configuration such that the configuration on the output tape is given by $\triangleright q_m f(w_1, \dots, w_\ell)$.

We assume that once T reaches an accepting configuration it stays indefinitely in that configuration (i.e., it loops). We further assume that T only reaches an accepting configuration when all its input words have the same size. Furthermore, when all inputs have the same size, T will reach an accepting configuration.

We say that T is a *linear space machine* when it reaches an accepting configuration on inputs of size n by using $\mathcal{O}(n)$ space on its work tape and additionally needs $\mathcal{O}(n^k)$ steps to do so. A *linear input-output function* is a function of the form $f = \bigcup_{n \geq 0} f_n : (\Sigma^n)^\ell \rightarrow \Sigma^n$. In other words, for every ℓ words of the same size n , f returns a word of size n . We say that a linear input-output function is a *linear space input-output function* if there exists a linear space machine T that for every $n \geq 0$, on input $w_1, \dots, w_\ell \in \Sigma^n$ the TM T has $f_n(w_1, \dots, w_\ell)$ on its the output tape when (necessarily) reaching an accepting configuration.

PROPOSITION 5.1. *Let $f = \bigcup_{n \geq 0} f_n : (\Sigma^n)^\ell \rightarrow \Sigma^n$ be a linear space input-output function computed by a linear space machine T with m states, ℓ input tapes, which consumes $\mathcal{O}(n)$ space and runs in $\mathcal{O}(n^{k-1})$ time on inputs of size n . There exists (i) a MATLANG schema $\mathcal{S} = (\mathcal{M}, \mathbf{size})$ where $\mathcal{M}$ consists matrix variables¹*

$$Q_1, \dots, Q_m, R_1, \dots, R_\ell, H_1, \dots, H_\ell, W_1, \dots, W_s, H_{W_1}, \dots, H_{W_s}, O, H_O, v_1, \dots, v_k$$

¹We also need a finite number of auxiliary variables, these will be specified in the proof.

with $\text{size}(V) = \alpha \times 1$ for all $V \in \mathcal{M}$; and (ii) a MATLANG expression e_f over $\mathcal{S}$ such that for the instance $\mathcal{I} = (\mathcal{D}, \text{mat})$ over $\mathcal{S}$ with $\mathcal{D}(\alpha) = n$ and

$$\text{mat}(R_i) = \text{vec}(w_i) \in \mathbb{R}^n,$$

for $i \in [\ell]$ and all other matrix variables instantiated with the zero vector in $\mathbb{R}^n$ for words $w_1, \dots, w_\ell \in \Sigma^n$ and such that $\text{vec}(w_i)$ is the $n \times 1$ -vector encoding the word w_i , we have that $\text{mat}(O) = \text{vec}(f_n(w_1, \dots, w_n)) \in \mathbb{R}^n$ after evaluating e_f on $\mathcal{I}$.

PROOF. The expression e_f we construct will simulate the TM T . To have some more control on the space and time consumption of T , let us first assume that n is large enough, say larger than $n \geq N$, such that T runs in sn space and $cn^{k-1} \leq n^k$ time for constants s and c . We deal with $n < N$ later on.

To simulate T we need to encode states, tapes and head positions. The matrix variables in $\mathcal{M}$ mentioned in the proposition will take these roles. More specifically, the variables $R_1, \dots, R_\ell$ will hold the input vectors, $W_1, \dots, W_s$ will hold the contents of the work tape, where s is the constant mentioned earlier, and O will hold the contents of the output tape. The vectors corresponding to the work and output tape are initially set to the zero vector. The vector for the input tape R_i is set to $\text{vec}(w_i)$, for $i \in [\ell]$.

With each tape we associate a matrix variable encoding the position of the head. More specifically, $H_1, \dots, H_\ell$ correspond to the input tape heads, $H_{W_1}, \dots, H_{W_s}$ are the heads for the work tape, and H_O is the head of the output tape. All these vectors are initialised with the zero vector. Later on, these vectors will be zero except for a single position, indicating the positions in the corresponding tapes the heads point to. For those positions j , $1 < j < n$, the head vectors will carry value 1. When $j = 1$ or n and when it concerns positions for the input tape, the head vectors can carry value 1 or 2. We need to treat these border cases separately because we only have n positions available to store the input words, whereas the actual input tapes consist of $n + 2$ symbols because of $\triangleright$ and $\triangleleft$. So when, for example, H_1 has a 1 in its first entry, we interpret it as the head is pointing to

the first symbol of the input word w_1 . When H_1 has a 2 in its first position, we interpret it as the head pointing to $\triangleright$. The end marker $\triangleleft$ is dealt with in the same way, by using value 1 or 2 in the last position of H_1 . We use this encoding for all input tapes, and also for the work tape W_1 and output tape O with the exception that no end marker $\triangleleft$ is present.

To encode the states, we use the variables $Q_1, \dots, Q_m$. We will ensure that when T is in state q_i when $\text{mat}(Q_i) = [1, 0, \dots, 0]^T \in \mathbb{R}^n$, otherwise $\text{mat}(Q_i)$ is the zero vector in $\mathbb{R}^n$.

Finally, the variables $v_1, \dots, v_k$ represent k canonical vectors which are used to iterate in for-loops. By iterating over them, we can perform n^k iterations, which suffices for simulating the $\mathcal{O}(n^{k-1})$ steps used by T to reach an accepting configuration.

With these matrix variables in place, we start by defining e_f . It will consist of two subexpressions $e_f^{\geq N}$, for dealing with $n \geq N$, and $e_f^{< N}$, for dealing with $n < N$. We explain the expression $e_f^{\geq N}$ first.

In our expressions we use subexpressions which we defined before in Chapter 3.2.2. These subexpressions require some auxiliary variables, as detailed below. As a consequence, e_f will be an expression defined over an extended schema $\mathcal{S}'$. Hence, the instance $\mathcal{I}$ in the statement of the Proposition is an instance $\mathcal{I}'$ of $\mathcal{S}'$ which coincides with $\mathcal{I}$ on $\mathcal{S}$ and in which the auxiliary matrix variables are all instantiated with zero vectors or matrices, depending on their size.

Now, we specify the finite auxiliary variables involved in the `for-MATLANG` expression. These arise when computing the following `for-MATLANG` expressions defined

- $e_{\text{Prev}}(z, Z, z', Z')$, and expression over auxiliary variables z, z', Z and Z' with $\text{size}(z) = \text{size}(z') = \text{size}(Z) = \alpha \times 1$ and $\text{size}(Z') = \alpha \times \alpha$. On input $\mathcal{I}'$ with $\text{mat}(z) = \text{mat}(z') = \text{mat}(Z)$ the zero column vector of dimension n , and $\text{mat}(Z')$ the zero $n \times n$ matrix, $\llbracket e_{\text{Prev}} \rrbracket(\mathcal{I}')$ returns the $n \times n$ matrix Prev such

that

$$\text{Prev} \cdot \mathbf{e}_i := \begin{cases} \mathbf{e}_{i-1} & \text{if } i > 1 \\ \mathbf{0} & \text{if } i = 1. \end{cases}$$

- $e_{\text{Next}}(z, Z, z', Z')$, and expression over auxiliary variables z, z', Z and Z' with $\text{size}(z) = \text{size}(z') = \text{size}(Z) = \alpha \times 1$ and $\text{size}(Z') = \alpha \times \alpha$. On input $\mathcal{I}'$ with $\text{mat}(z) = \text{mat}(z') = \text{mat}(Z)$ the zero column vector of dimension n , and $\text{mat}(Z')$ the zero $n \times n$ matrix, $\llbracket e_{\text{Next}} \rrbracket(\mathcal{I}')$ returns the $n \times n$ matrix Next such that

$$\text{Next} \cdot \mathbf{e}_i := \begin{cases} \mathbf{e}_{i+1} & \text{if } i < n \\ \mathbf{0} & \text{if } i = n. \end{cases}$$

- $\min(v, z, Z, z', Z)$ with auxiliary variables z, z', Z and Z' as before, and v is one of the (vector) variables in $\mathcal{M}$. For an $n \times 1$ vector $\mathbf{v}$, on input $\mathcal{I}'[v \leftarrow \mathbf{v}]$

$$\llbracket \min \rrbracket(\mathcal{I}'[v \leftarrow \mathbf{v}]) := \begin{cases} 1 & \text{if } \mathbf{v} = \mathbf{e}_1 \\ 0 & \text{otherwise.} \end{cases}$$

- $\max(v, z, Z, z', Z)$ with auxiliary variables z, z', Z and Z' as before, and v is one of the (vector) variables in $\mathcal{M}$. For an $n \times 1$ vector $\mathbf{v}$, on input $\mathcal{I}'[v \leftarrow \mathbf{v}]$

$$\llbracket \max \rrbracket(\mathcal{I}'[v \leftarrow \mathbf{v}]) := \begin{cases} 1 & \text{if } \mathbf{v} = \mathbf{e}_n \\ 0 & \text{otherwise.} \end{cases}$$

- $e_{\min}(z, Z, z', Z', z'', Z'')$, an expressions with auxiliary variables z, z', z'', Z, Z' and Z'' with $\text{size}(z) = \text{size}(z') = \text{size}(z'') = \text{size}(Z) = \text{size}(Z'') = \alpha \times 1$ and $\text{size}(Z') = \alpha \times \alpha$. On input $\mathcal{I}'$ with matrix variables instantiated with zero vectors (or matrix for Z'), $\llbracket e_{\min} \rrbracket(\mathcal{I}') = \mathbf{e}_1$.
- $e_{\max}(z, Z, z', Z', z'', Z'')$, an expressions with auxiliary variables z, z', z'', Z, Z' and Z'' with $\text{size}(z) = \text{size}(z') = \text{size}(z'') = \text{size}(Z) = \text{size}(Z'') = \alpha \times 1$ and $\text{size}(Z') = \alpha \times \alpha$. On input $\mathcal{I}'$ with matrix variables instantiated with zero vectors (or matrix for Z'), $\llbracket e_{\max} \rrbracket(\mathcal{I}') = \mathbf{e}_n$.

We thus see that we only need z, z', z'', Z, Z', Z'' as auxiliary variables and these can be re-used whenever e_f calls these functions. From now on, we omit the auxiliary variables from the description of e_f .

Let us first define $e_f^{\geq N}$. Since we want to simulate T we need to be able to check which transitions of T can be applied based on a current configuration. More precisely, suppose that we want to check whether $\delta(q_i, (a_1, \dots, a_\ell, b, c))$ is applicable, then we need to check whether T is in state q_i , we can do this by checking $\min(Q_i)$, and whether the heads on the tapes read symbols $a_1, \dots, a_\ell, b, c$. We check the latter by the following expressions. For the input tape R_i we define

$$\text{test_inp}_b^i := \begin{cases} (1 - \min(1/2 \cdot H_i)) \cdot (1 - \max(1/2 \cdot H_i)) \cdot (1 - R_i^T \cdot H_i) & \text{if } b = 0 \\ (1 - \min(1/2 \cdot H_i)) \cdot (1 - \max(1/2 \cdot H_i)) \cdot (R_i^T \cdot H_i) & \text{if } b = 1 \\ \min(1/2 \cdot H_i) & \text{if } b = \triangleright \\ \max(1/2 \cdot H_i) & \text{if } b = \triangleleft, \end{cases}$$

which returns 1 if and only if either $b \in \{0, 1\}$ is the value in $\text{mat}(R_i)$ at the position encoded by $\text{mat}(H_i)$, or when $b = \triangleright$ and $\text{mat}(H_i)$ is the vector $(2, 0, \dots, 0) \in \mathbb{R}^n$, or when $b = \triangleleft$ and $\text{mat}(H_i)$ is the vector $(0, 0, \dots, 2) \in \mathbb{R}^n$. Similarly, for the output tape we define

$$\text{test_out}_b := \begin{cases} (1 - \min(1/2 \cdot H_O)) \cdot (1 - O^T \cdot H_O) & \text{if } b = 0 \\ (1 - \min(1/2 \cdot H_O)) \cdot (O^T \cdot H_O) & \text{if } b = 1 \\ \min(1/2 \cdot H_O) & \text{if } b = \triangleright, \end{cases}$$

and for the work tapes $W_1, \dots, W_s$ we define

$$\text{test_work}_b^i := \begin{cases} (1 - \min(1/2 \cdot H_{W_i})) \cdot (1 - W_i^T \cdot H_{W_i}) & \text{if } b = 0 \\ (1 - \min(1/2 \cdot H_{W_i})) \cdot (W_i^T \cdot H_{W_i}) & \text{if } b = 1 \\ \min(1/2 \cdot H_{W_i}) & \text{if } b = \triangleright \text{ and } i = 1. \end{cases}$$

We then combine all these expressions into a single expression for $q_i \in Q, a_1, \dots, a_\ell, b, c \in \Gamma$:

$$\text{isconf}_{q_i, a_1, \dots, a_\ell, b, c} := \min(Q_i) \cdot \left(\prod_{j=1}^{\ell} \text{test_inp}_{a_j}^j \right) \cdot \left(\sum_{j=1}^s \text{test_work}_b^j \right) \cdot \text{test_out}_c.$$

This expression will return 1 if and only if the vectors representing the tapes, head positions and states are such that Q_i is the first canonical vector (and thus T is in state q_i), the heads point to entries in the tape vectors storing the symbols $a_1, \dots, a_\ell, b, c$ or they point to the first (or last for input tapes) positions but have value 2 (when the symbols are $\triangleright$ or $\triangleleft$).

To ensure that at the beginning of the simulation of T by $e_f^{\geq N}$ we correctly encode that we are in the initial configuration, we thus need to initialise all vectors

$$\text{mat}(H_1), \text{mat}(H_2), \dots, \text{mat}(H_\ell), \text{mat}(H_{W_1}), \text{mat}(H_O)$$

with the vector $(2, 0, 0, \dots, 0) \in \mathbb{R}$ since all heads read the symbol $\triangleright$. Similarly, we have to initialise Q_1 with the first canonical vector since T is in state q_0 .

We furthermore need to be able to correctly adjust head positions. We do this by means of the predecessor and successor expressions described above. A consequence of our encoding is that we need to treat the border cases (corresponding to $\triangleright$ and $\triangleleft$) differently. More specifically, for the input tapes R_i and heads H_i we define

$$\text{move_inp}_d^i := \begin{cases} 2 \times \min(H_i) \times H_i + 1/2 \times \max(1/2 \times H_i) \times H_i + \\ \quad (1 - \min(H_i))(1 - \max(1/2 \times H_i)) \times e_{\text{Prev}} \cdot H_i & \text{if } d = \leftarrow \\ 2 \times \max(H_i) \times H_i + 1/2 \times \min(1/2 \times H_i) \times H_i + \\ \quad (1 - \min(1/2 \times H_i))(1 - \max(H_i)) \times e_{\text{Next}} \cdot H_i & \text{if } d = \rightarrow \\ H_i & \text{if } d = \sqcup. \end{cases}$$

In other words, we shift to the previous (or next) canonical vector when d is $\leftarrow$ or $\rightarrow$, respectively, unless we need to move to or from the position that will hold $\triangleright$ or $\triangleleft$. In those case we readjust $\text{mat}(H_i)$ (which will either $(1, 0, \dots, 0)$, $(2, 0, \dots, 0)$, $(0, \dots, 0, 1)$ or $(0, \dots, 0, 2)$) by either dividing or multiplying with 2. In this way we can correctly infer whether or not the head points to the begin and end markers. For the output tape we proceed in a similar way, but only taking into account the begin marker and recall that we do not have moves to the left:

$$\text{move_outp}_d := \begin{cases} 1/2 \times \min(1/2 \times H_O) \times H_O + \\ \quad (1 - \min(1/2 \times H_O)) \times e_{\text{Next}} \times H_O & \text{if } d = \rightarrow \\ H_O & \text{if } d = \sqcup. \end{cases}$$

Since we represent the work tape by s vectors $W_1, \dots, W_s$ we need to ensure that only one of the head vectors H_{W_i} has a non-zero value and that by moving left or right, we need to appropriately update the right head vector. We do this as follows. We first consider the work tapes W_i for $i \neq 1, s$ and define

$$\text{move_work}_d^i := \begin{cases} -\min(H_{W_i}) \times H_{W_i} + (1 - \min(H_{W_i})) \times e_{\text{Prev}} \cdot H_{W_i} + \\ \quad \min(H_{W_{i+1}}) \times e_{\text{max}} & \text{if } d = \leftarrow \\ -\max(H_{W_i}) \times H_{W_i} + (1 - \max(H_{W_i})) \times e_{\text{Next}} \cdot H_{W_i} + \\ \quad \max(H_{W_{i-1}}) \times e_{\text{min}} & \text{if } d = \rightarrow \\ H_{W_i} & \text{if } d = \sqcup. \end{cases}$$

In other words, we set the H_{W_i} to zero when a move brings us to either W_{i-1} or W_{i+1} , we move the successor or predecessor when staying within W_i , or initialise H_{W_i} with the first or last canonical vector when moving from W_{i-1} to W_i (right move) or from W_{i+1} to W_i (left move). For $i = s$ we can ignore the parts in the previous expression that involve

W_{s+1} (which does not exist):

$$\text{move_work}_d^s := \begin{cases} -\min(H_{W_s}) \times H_{W_i} + (1 - \min(H_{W_s})) \times e_{\text{Prev}} \cdot H_{W_s} & \text{if } d = \leftarrow \\ -\max(H_{W_s}) \times H_{W_s} + (1 - \max(H_{W_s})) \times e_{\text{Next}} \cdot H_{W_s} + \\ \quad \max(H_{W_{s-1}}) \times e_{\min} & \text{if } d = \rightarrow \\ H_{W_s} & \text{if } d = \sqcup. \end{cases}$$

For $i = 1$, we can ignore the part involving W_0 (which does not exist) but have to take $\triangleright$ into account:

$$\text{move_work}_d^1 := \begin{cases} 2 \times \min(H_{W_1}) \times H_{W_i} + (1 - \min(H_{W_1})) \times e_{\text{Prev}} \cdot H_{W_1} + \\ \quad \min(H_{W_2}) \times e_{\max} & \text{if } d = \leftarrow \\ 1/2 \times \min(1/2 \times H_{W_1}) \times H_{W_1} + \\ \quad (1 - \max(1/2 \times H_{W_1})) \times e_{\text{Next}} \cdot H_{W_1} & \text{if } d = \rightarrow \\ H_{W_1} & \text{if } d = \sqcup. \end{cases}$$

A final ingredient for defining $e_f^{\geq N}$ are expressions which update the work and output tape. To define these expression, we need the position and symbol to put on the tape. For the output tape we define

$$\text{write_outp}_b := \begin{cases} \min(1/2 \times H_O) \times O & \text{if } b = \triangleright \\ (1 - \min(1/2 \times H_O)) \times \\ \quad ((1 - O^T \cdot H_O) \times O + (O^T \cdot H_O) \times (O - H_O)) & \text{if } b = 0 \\ (1 - \min(1/2 \times H_O)) \times \\ \quad ((1 - O^T \cdot H_O) \times (O + H_O) + (O^T \cdot H_O) \times O) & \text{if } b = 1 \end{cases}$$

and similarly for the work tapes $i \neq 1$:

$$\text{write_work}_b^i := \begin{cases} W_i & \text{if } b = \triangleright \\ (1 - W_i^T \cdot H_{W_i}) \times W_i + (W_i^T \cdot H_{W_i}) \times (W_i - H_{W_i}) & \text{if } b = 0 \\ (1 - W_i^T \cdot H_{W_i}) \times (W_i + H_{W_i}) + (W_i^T \cdot H_{W_i}) \times W_i & \text{if } b = 1, \end{cases}$$

and for W_1 we have to take care again of the begin marker:

$$\text{write_work}_b^1 := \begin{cases} \min(1/2 \times H_{W_1}) \times W_1 & \text{if } b = \triangleright \\ (1 - \min(1/2 \times H_{W_1}) \times \\ \quad \left((1 - W_1^T \cdot H_{W_1}) \times W_1 + \right. \\ \quad \quad \left. (W_1^T \cdot H_{W_1}) \times (W_1 - H_{W_1}) \right) & \text{if } b = 0 \\ (1 - \min(1/2 \times H_{W_1}) \times \\ \quad \left((1 - W_1^T \cdot H_{W_1}) \times (W_1 + H_{W_1}) + \right. \\ \quad \quad \left. (W_1^T \cdot H_{W_1}) \times W_1 \right) & \text{if } b = 1. \end{cases}$$

We are now finally ready to define $e_f^{\geq N}$:

$$e_f^{\geq N} := \text{for } v_1, \dots, v_k, Q_1, \dots, Q_m, H_1, \dots, H_\ell, W_1, \dots, W_s, H_{W_1}, \dots, H_{W_s}, O, H_O. \\ (e_{Q_1}, \dots, e_{Q_m}, e_{H_1}, \dots, e_{H_\ell}, e_{W_1}, \dots, e_{W_s}, e_{H_{W_1}}, \dots, e_{H_{W_s}}, e_O, e_{H_O}).$$

We refer to Chapter 3.3.1 for the definition of this form of the for-loop. The expressions used are (we use $\star$ below to mark irrelevant information in the transitions):

$$\begin{aligned}
e_{Q_1} &:= \left(\prod_{j=1}^k \min(v_i) \right) \times e_{\min} + \sum_{\substack{(q_i, a_1, \dots, a_\ell, b, c) \\ \Delta(q_i, a_1, \dots, a_\ell, b, c) = (q_1, \star)}} \text{isconf}_{q_i, a_1, \dots, a_\ell, b, c} \times e_{\min} \\
e_{Q_j} &:= \sum_{\substack{(q_i, a_1, \dots, a_\ell, b, c) \\ \Delta(q_i, a_1, \dots, a_\ell, b, c) = (q_j, \star)}} \text{isconf}_{q_i, a_1, \dots, a_\ell, b, c} \times e_{\min} \quad \text{for } j \neq 1 \\
e_{H_i} &:= 2 \left(\prod_{j=1}^k \min(v_i) \right) \times e_{\min} + \sum_{\substack{(q, a_1, \dots, a_\ell, b, d) \\ \Delta(q, a_1, \dots, a_\ell, b, c) = (\star, d_i, \star)}} \text{isconf}_{q, a_1, \dots, a_\ell, b, c} \times \text{move_inp}_{d_i}^i \\
e_{H_{W_i}} &:= 2 \left(\prod_{j=1}^k \min(v_i) \right) \times e_{\min} + \sum_{\substack{(q, a_1, \dots, a_\ell, b, d) \\ \Delta(q, a_1, \dots, a_\ell, b, c) = (\star, d_{\ell+1}, \star)}} \text{isconf}_{q, a_1, \dots, a_\ell, b, c} \times \text{move_work}_{d_{\ell+1}}^i \\
e_{H_O} &:= 2 \left(\prod_{j=1}^k \min(v_i) \right) \times e_{\min} + \sum_{\substack{(q, a_1, \dots, a_\ell, b, d) \\ \Delta(q, a_1, \dots, a_\ell, b, c) = (\star, d_{\ell+2}, \star)}} \text{isconf}_{q, a_1, \dots, a_\ell, b, c} \times \text{move_outp}_{d_{\ell+2}} \\
e_{W_i} &:= \sum_{\substack{(q, a_1, \dots, a_\ell, b, d) \\ \Delta(q, a_1, \dots, a_\ell, b, c) = (\star, b', c', \star)}} \text{isconf}_{q, a_1, \dots, a_\ell, b, c} \times \text{write_work}_{b'}^i \\
e_O &:= \sum_{\substack{(q, a_1, \dots, a_\ell, b, d) \\ \Delta(q, a_1, \dots, a_\ell, b, c) = (\star, b', c', \star)}} \text{isconf}_{q, a_1, \dots, a_\ell, b, c} \times \text{write_outp}_{c'}.
\end{aligned}$$

The correctness of $e_f^{\geq N}$ should be clear from the construction (one can formally verify this by induction on the number of iterations). We use $\prod_{j=1}^k \min(v_i)$ to initialize with $e_{\min}$, by checking that the v_i 's are all in its first assignment. We next explain how the border cases $n < N$ can be dealt with. For each $n < N$ and every possible input words $w_1, \dots, w_\ell$ of size n , we define a `for-MATLANG` expression which checks whether $\text{mat}(R_i) = \text{vec}(w_i)$ for each $i \in [\ell]$. This can be easily done since n can be regarded as a constant. For example, to check whether $\text{mat}(R_i) = [0, 1, 1]^T$ we simply write

$$(1 - R_i^T \cdot e_{\min}) \times (R_i^T \cdot e_{\text{Next}} \cdot e_{\min}) \times (1 - R_i^T \cdot e_{\text{Next}} \cdot e_{\text{Next}} \cdot e_{\min}) \times (1 - e_1(R_i)^T \cdot e_{\text{Next}} \cdot e_{\text{Next}} \cdot e_{\text{Next}} \cdot e_{\min})$$

which will evaluate to 1 if and only if $\text{mat}(R_i) = [0, 1, 1]^T$. We note that the final factor is in place to check that the dimension of $\text{mat}(R_i)$ is three. We denote by $e_{n,w}^i$ the expression which evaluates to 1 if and only if $\text{mat}(R_i) = \text{vec}(w)$ for $|w| = n$. We can similarly write any word w of fixed size in the matrix variable O . For example, suppose that $w = 101$ then we write

$$O + e_{\min} + e_{\text{Next}} \cdot e_{\text{Next}} \cdot e_{\min}.$$

We write $e_{n,w}$ be the expression which write w of size $|w| = n$ in matrix variable O . Then, the expressions

$$e_{n,w_1,\dots,w_n,w} := e_{n,w_1}^1 \cdot \dots \cdot e_{n,w_\ell}^\ell \cdot e_{n,w}$$

will write w in O if and only if $\text{mat}(R_i) = \text{vec}(w_i)$ for $i \in [\ell]$. We now simply take the disjunction over all words $w_1, \dots, w_\ell \in \Sigma^n$ and $w = f_n(w_1, \dots, w_\ell) \in \Sigma^n$:

$$e_n := \sum_{w_1, \dots, w_\ell \in \Sigma^n} e_{n,w_1, \dots, w_\ell, f_n(w_1, \dots, w_\ell)},$$

which correctly evaluates f_n . We next take a further disjunction by letting ranging from $n = 0, \dots, N - 1$:

$$e_f^{<N} := \sum_{n=0}^{N-1} e_n$$

Since every possible input is covered and only a unique expression $e_{n,w_1, \dots, w_\ell, f_n(w_1, \dots, w_\ell)}$ will be triggered $e_f^{<N}$ will correctly evaluate f on inputs smaller than N .

Our final expression e_f is now given by

$$e_f := e_f^{<N} + \text{dim_is_greater_than}_N \times e_f^{\geq N}$$

where $\text{dim_is_greater_than}_N$ is the expression $e_1(R_i)^T \cdot \underbrace{e_{\text{Next}} \cdot \dots \cdot e_{\text{Next}}}_{N \text{ times}}$ which will evaluate to 1 if and only if the input dimension is larger or equal than N . $\square$

Note that we need the computation to use $\mathcal{O}(n)$ space since we store the contents of the work tape in n dimensional vectors. If we allow quadratic work space, these vectors would have n^2 size, thus yielding dimension mismatch in matrix operations and undefined

for-MATLANG expressions. The same reasoning applies to allowing quadratic output space.

5.2. From arithmetic circuits to for-MATLANG

Let us first recall the definition of arithmetic circuits. An *arithmetic circuit* Φ over a set $X = \{x_1, \dots, x_n\}$ of input variables is a directed acyclic labeled graph. The vertices of Φ are called *gates* and denoted by $g_1, \dots, g_m$; the edges in Φ are called *wires*. The children of a gate g correspond to all gates g' such that (g, g') is an edge. The parents of g correspond to all gates g' such that (g', g) is an edge. The *in-degree*, or a *fan-in*, of a gate g refers to its number of children, and the *out-degree* to its number of parents. We will not assume any restriction on the in-degree of a gate, and will thus consider circuits with unbounded fan-in. Gates with in-degree 0 are called *input gates* and are labeled by either a variable in X or a constant 0 or 1. All other gates are labeled by either $+$ or $\times$, and are referred to as *sum gates* or *product gates*, respectively. Gates with out-degree 0 are called *output gates*. When talking about arithmetic circuits, one usually focuses on circuits with n input gates and a single output gate.

The *size* of Φ , denoted by $|\Phi|$, is its number of gates and wires. The *depth* of Φ , denoted by $\text{depth}(\Phi)$, is the length of the longest directed path from any of its output gates to any of the input gates. The *degree* of a gate is defined inductively: an input gate has degree 1, a sum gate has a degree equal to the maximum of degrees of its children, and a product gate has a degree equal to the sum of the degrees of its children. When Φ has a single output gate, the *degree* of Φ , denoted by $\text{degree}(\Phi)$, is defined as the degree of its output gate. If Φ has a single output gate and its input gates take values from $\mathbb{R}$, then Φ corresponds to a polynomial in $\mathbb{R}[X]$ in a natural way. In this case, the degree of Φ equals the degree of the polynomial corresponding to Φ .

In order to handle inputs of different sizes, we use the notion of uniform circuit families. An *arithmetic circuit family* is a set of arithmetic circuits $\{\Phi_n \mid n = 1, 2, \dots\}$ where

Φ_n has n input variables and a single output gate. An arithmetic circuit family is *uniform* if there exists a LOGSPACE-Turing machine, which on input 1^n , returns an encoding of the arithmetic circuit Φ_n for each n . We observe that uniform arithmetic circuit families are necessarily of polynomial size. Another important parameter is the circuit depth. A circuit family is of logarithmic depth, whenever $\text{depth}(\Phi_n) \in \mathcal{O}(\log n)$. We now show that `for-MATLANG` subsumes uniform arithmetic circuit families that are of logarithmic depth.

Theorem 5.1. *For any uniform arithmetic circuit family $\{\Phi_n \mid n = 1, 2, \dots\}$ of logarithmic depth there is a `for-MATLANG` schema $\mathcal{S}$ and an expression e_Φ using a matrix variable v , with $\text{type}_\mathcal{S}(v) = (\alpha, 1)$ and $\text{type}_\mathcal{S}(e) = (1, 1)$, such that for any input values $a_1, \dots, a_n$:*

- *If $\mathcal{I} = (\mathcal{D}, \text{mat})$ is a `MATLANG` instance such that $\mathcal{D}(\alpha) = n$ and $\text{mat}(v) = [a_1 \dots a_n]^T$.*
- *Then $\llbracket e_\Phi \rrbracket(\mathcal{I}) = \Phi_n(a_1, \dots, a_n)$.*

It is important to note that the expression e_Φ does not change depending on the input size, meaning that it is uniform in the same sense as the circuit family being generated by a single Turing machine. The different input sizes for a `for-MATLANG` instance are handled by the typing mechanism of the language.

The proof of this Theorem, which is the deepest technical result of the thesis, depends crucially on two facts: (i) that any polynomial time Turing machine working within linear space and producing linear size output, can be simulated via a `for-MATLANG` expression; and (ii) that evaluating an arithmetic circuit Φ_n can be done using two stacks of depth n .

Evaluating Φ_n on input $(a_1, \dots, a_n)$ can be done in a depth-first manner by maintaining two stacks: the gates-stack that tracks the current gate being evaluated, and the values-stack that stores the value that is being computed for this gate. The idea behind having two stacks is that whenever the number of items on the gates-stack is higher by one than the number of items on the values-stack, we know that we are processing a fresh gate, and

we have to initialize its current value (to 0 if it is a sum gate, and to 1 if it is a product gate), and push it to the values-stack. We then proceed by processing the children of the head of the gates-stack one by one, and aggregate the results using sum if we are working with a sum gate, and by using product otherwise.

In order to access the information about the gate we are processing (such as whether it is a sum or a product gate, the list of its children, etc.) we use the uniformity of our circuit family. Namely, we know that we can generate the circuit Φ_n with a LOGSPACE-Turing machine M_Φ by running it on the input 1^n . Using this machine, we can in fact compute all the information needed to run the two-stack algorithms described above. For instance, we can construct a LOGSPACE machine that checks, given two gates g_1 and g_2 , whether g_2 is a child of g_1 . Similarly, we can construct a machine that, given g_1 and g_2 tells us whether g_2 is the final child of g_1 , or the one that produces the following child of g_1 (according to the ordering given by the machine M_Φ). Defining these machines based of M_Φ is similar to the algorithm for the composition of two LOGSPACE transducers, and is commonly used to evaluate arithmetic circuits (Allender, 2004).

To simulate the circuit evaluation algorithm that uses two stacks, in `for-MATLANG` we can use a binary matrix of size $n \times n$, where n is the number of inputs. The idea here is that the gates-stack corresponds to the first $n - 3$ columns of the matrix, with each gate being encoded as a binary number in positions $1, \dots, n - 3$ of a row. The remaining three columns are reserved for the values-stack, the number of elements on the gates stack, and the number of elements on the values stack, respectively. The number of elements is encoded as a canonical vector of size n . Here we crucially depend on the fact that the circuit is of logarithmic depth, and therefore the size of the two stacks is bounded by n (apart from the portion before the asymptotic bound kicks-in, which can be hard-coded into the expression e_Φ). Similarly, given that the circuits are of polynomial size, we can assume that gate ids can be encoded into $n - 3$ bits.

This matrix is then updated in the same way as the two-stack algorithm. It processes gates one by one, and using the successor relation for canonical vectors determines

whether we have more elements on the gates stack. In this case, a new value is added to the values stack (0 if the gate is a sum gate, and 1 otherwise), and the process continues. Information about the next child, last child, or input value, are obtained using the expression which simulates the Turing machine generating this data about the circuit (the machines used never produce an output longer than their input). Given that the size of the circuit is polynomial, say n^k , we can initialize the matrix with the output gate only, and run the simulation of the two-stack algorithm for n^k steps (by iterating k times over size n canonical vectors). After this, the value in position $(1, n - 2)$ (the top of the values stack) holds the final results.

Proof of theorem 5.1. Let Φ_n be a circuit with n input gates and such that it can be computed by a L -uniform arithmetic circuit of log-depth. Each gate of the circuit that encodes f has an $\text{id} \in \{0, 1\}^n$. From now on, when we write g for a gate of the circuit, we mean the id encoding g . Let n^k be a polynomial such that the number of wires $W(n) \leq n^k$ for n big enough. Further, we assume that $2W(n) \leq n^k$. We need this because the for-matlang simulation of the circuit is in a depth first search way, so $2W(n)$ wires will be traversed. Then we have that:

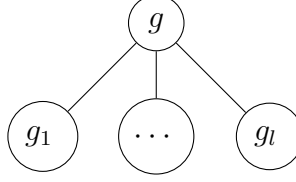
- the number of gates is bounded by n^k .
- we need at most $k \log(n)$ bits to store the id of a gate.
- the depth of the circuit is at most $k' \log(n)$ for some k' .

So, let n_0 and k such that $\forall n \geq n_0$:

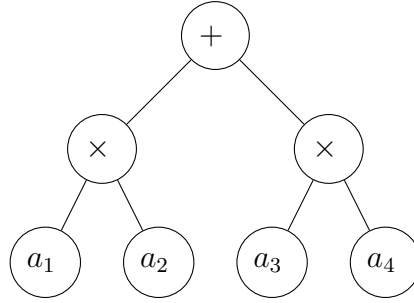
$$\left. \begin{array}{l} 2W(n) \leq n^k \\ k \lceil \log(n) \rceil \leq n - 3 \\ k' \lceil \log(n) \rceil \leq n \end{array} \right\} (\star)$$

We know n_0 and k exist. Let $n \geq n_0$. Towards the end, we will deal with the case when $n < n_0$.

Let g be a gate. The children of g are denoted by $g_1, \dots, g_l$.



For example, a circuit that encodes the function $f(a_1, a_2, a_3, a_4) = a_1a_2 + a_3a_4$ is



We can simulate the polynomial $x^2 + xy$ by doing $f(A)$ where $A = [x \ x \ x \ y]^T$. The main idea is to traverse the circuit top down in a depth first search way and store visited gates in a stack and its corresponding current values in another stack, and aggregate in the iterations according to the gate type.

For a stack S , the operations are standard:

- $S.\text{push}(s)$: pushes s into S .
- $S.\text{pop}$: pops the top element.
- $S.\text{size}$: the length of the stack.
- $S.\text{top}$: the top element in the stack.

For the pseudo-code, $\mathcal{G}$ and $\mathcal{V}$ denote stacks of gates and values, respectively. The property that holds during the simulation is that the value in $\mathcal{V}[i]$ is the value that $\mathcal{G}[i]$ currently outputs. The algorithm ends with $\mathcal{G} = [g_{\text{root}}]$ and $\mathcal{V} = [v_{\text{root}}]$ after traversing the circuit, and returns v_{root}

During the evaluation algorithm there will be two possible configurations of $\mathcal{G}$ and $\mathcal{V}$.

- (i) $\mathcal{G}.size = \mathcal{V}.size + 1$: this means that $\mathcal{G}.top$ is a gate that we visit for the first time and we need to initialize its value.
- (ii) $\mathcal{G}.size = \mathcal{V}.size$: here $\mathcal{V}.top$ is the value of evaluating the circuit in gate $\mathcal{G}.top$. Therefore, we need to aggregate the value $\mathcal{V}.top$ to the parent gate of g .

We assume the circuit has input gates, $+$, $\times$ -gates and allow constant 1-gate.

The idea is to traverse the circuit top down in a depth first search way. For example, in the circuit $f(a_1, a_2, a_3, a_4) = a_1a_2 + a_3a_4$ above, we would initialize the output gate value as 0 because it is a $+$ gate, so $\mathcal{G} = \{+\}$, $\mathcal{V} = \{0\}$. Then stack the left $\times$ gate to $\mathcal{G}$, stack its initial value (i.e. 1) to $\mathcal{V}$. Now stack a_1 to $\mathcal{G}$ and its value (i.e. a_1) to $\mathcal{V}$. Since we are on an input gate we pop the gate and value pair off of $\mathcal{G}$ and $\mathcal{V}$ respectively, aggregate a_1 to $\mathcal{V}.top$ and continue by stacking the a_2 gate to $\mathcal{G}$. We pop a_2 off of $\mathcal{V}$ (and its gate off of $\mathcal{G}$) and aggregate its value to $\mathcal{V}.top$. We pop and aggregate the value of the left $\times$ gate to $\mathcal{V}.top$ (the root value). Then continue with the right $\times$ gate branch similarly.

For the pseudo-code, we supply ourselves with the following functions:

- `isplus( $g$ )`: true if and only if g is a $+$ -gate.
- `isprod( $g$ )`: true if and only if g is a $\times$ -gate.
- `isone( $g$ )`: true if and only if g is a 1-gate.
- `isinput( $g$ )`: true if and only if g is an input gate.
- `getfirst( $g$ )`: outputs the first child of g .
- `getinput( $g$ )`: outputs $A[i]$ when g is the i -th input.
- `not_last( $g_1, g_2$ )`: true if and only if g_2 is not the last child gate of g_1 .
- `next_gate( $g_1, g_2$ )`: outputs the next child gate of g_1 after g_2 .
- `getroot()`: outputs the root gate of the circuit.

The corresponding $\{0, 1\}^n \rightarrow \{0, 1\}^n$ functions are:

- `isplus(g)`: 1 if and only if g is a $+$ -gate.
- `isprod(g)`: 1 if and only if g is a $\times$ -gate.
- `isone(g)`: 1 if and only if g is a 1-gate.
- `isinput(g)`: 1 if and only if g is an input gate.
- `getfirst(g)`: outputs the `id` of the first child of g .
- `getinput(g)`: outputs canonical vector b_i , where the i -th input gate of Φ_n is encoded by g .
- `not_last(g1, g2)`: 1 if and only if g_2 is not the last child gate of g_1 .
- `next_gate(g1, g2)`: outputs the `id` of the next child gate of g_1 after g_2 .
- `getroot()`: outputs the `id` of the root gate of the circuit.

The previous functions are all definable by an L -transducer and can be defined from the L -transducer of f . Then, by proposition 5.1, for each of these functions there is a `for-MATLANG` expression that simulates them.

Now, we give the pseudo-code of the top-down evaluation. We define the functions *Initialize* (algorithm 1), *Aggregate* (algorithm 2) and *Evaluate* (algorithm 3). The main algorithm is *Evaluate*.

Algorithm 1 Initialize (pseudo-code)

```

1: function INITIALIZE( $\mathcal{G}, \mathcal{V}, A$ )                                ▷ The stacks and input. Here,
    $\mathcal{G}.size = \mathcal{V}.size + 1$ 
2:   if isplus( $\mathcal{G}.top$ ) then
3:      $\mathcal{V}.push(0)$ 
4:      $\mathcal{G}.push(\text{getfirst}(\mathcal{G}.top))$ 
5:   else if isprod( $\mathcal{G}.top$ ) then
6:      $\mathcal{V}.push(1)$ 
7:      $\mathcal{G}.push(\text{getfirst}(\mathcal{G}.top))$ 
8:   else if isone( $\mathcal{G}.top$ ) then
9:      $\mathcal{V}.push(1)$ 
10:  else if isinput( $\mathcal{G}.top$ ) then
11:     $\mathcal{V}.push(A[\text{getinput}(\mathcal{G}.top)])$ 
12:  return  $\mathcal{G}, \mathcal{V}$ 

```

Algorithm 2 Aggregate (pseudo-code)

```
1: function AGGREGATE( $\mathcal{G}, \mathcal{V}$ ) ▷ Here,  $\mathcal{G}.size = \mathcal{V}.size$ 
2:    $g = \mathcal{G}.pop$ 
3:    $v = \mathcal{V}.pop$ 
4:   if isplus( $\mathcal{G}.top$ ) then
5:      $\mathcal{V}.top = \mathcal{V}.top + v$ 
6:   else if isprod( $\mathcal{G}.top$ ) then
7:      $\mathcal{V}.top = \mathcal{V}.top \cdot v$ 
8:   if not_last( $\mathcal{G}.top, g$ ) then
9:      $\mathcal{G}.push(next\_gate(\mathcal{G}.top, g))$ 
10:  return  $\mathcal{G}, \mathcal{V}$ 
```

Algorithm 3 Evaluate (pseudo-code)

```
1: function EVALUATE( $A$ ) ▷ Input  $n \times 1$  vector  $A$ . Here,  $\mathcal{G}$  and  $\mathcal{V}$  are empty
2:    $\mathcal{G}.push(getroot())$ 
3:   while  $\mathcal{G}.size \neq 1$  or  $\mathcal{V}.size \neq 1$  do
4:     if  $\mathcal{G}.size \neq \mathcal{V}.size$  then
5:        $(\mathcal{G}, \mathcal{V}) := \text{Initialize}(\mathcal{G}, \mathcal{V}, A)$ 
6:     else
7:        $(\mathcal{G}, \mathcal{V}) := \text{Aggregate}(\mathcal{G}, \mathcal{V})$ 
8:   return  $\mathcal{V}.top$ 
```

The *Evaluate* algorithm gives us the output of the circuit. Note that after each iteration it either holds that $\mathcal{G}.size = \mathcal{V}.size + 1$ or $\mathcal{G}.size = \mathcal{V}.size$. Furthermore, when we start we have $\mathcal{G}.size = 1$ and $\mathcal{V}.size = 0$. The condition $\mathcal{G}.size = 1$ and $\mathcal{V}.size = 1$ holds only when we have traversed all the circuit, and the value in $\mathcal{V}.top$ is the value that the root of the circuit outputs after its computation.

Next, we show how to encode this algorithm in `for-MATLANG`.

Let $n_0 \in \mathbb{N}$ be big enough for $(\star)$ to hold and let $n \geq k$. Hence, the number of gates (values) is bounded by n^k and we need $k \log(n)$ bits to encode the id of each gate.

To simulate the two stacks $\mathcal{G}$ and $\mathcal{V}$ we keep a matrix X of dimensions $n \times n$.

- Column n will store a canonical vector that marks the top of stack V (values).
- Column $n - 1$ will store a canonical vector that marks the top of stack G (gates).

- Column $n - 2$ is the stack of values where $X[1, n - 2]$ is the bottom of the stack.
- Columns 1 to $n - 3$ are the stack of gates.

If we have j gates in the stack and currently $\mathcal{G}.size = \mathcal{V}.size$ then X would look like:

$$X = \begin{bmatrix} \text{id}_1 & v_1 & 0 & 0 \\ \text{id}_2 & v_2 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ \text{id}_j & v_j & 1 & 1 \\ 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

Since $n \geq n_0$, $(\star)$ holds and thus we never use more than $n - 3$ bits to encode an id . Also, $j \leq n$ given that we never keep more gates than the depth of the tree. As a consequence, we never keep more than n values either.

An important detail is that the ids of the gates are encoded as id_r000 for it to have dimension n , where id_r is the corresponding binary number in reverse.

We make a series of definitions to make the notation more clear. Refer to Chapter 3.2.2 for more information about these expressions.

Let b_i be the i -th canonical vector. Next and Prev denote the successor and predecessor matrices respectively, such that

$$\text{Next} \cdot b_i = \begin{cases} b_{i+1} & \text{if } i \leq n \\ 0 & \text{otherwise} \end{cases}$$

$$\text{Prev} \cdot b_i = \begin{cases} b_{i-1} & \text{if } i \geq n \\ 0 & \text{otherwise} \end{cases}$$

We write expressions $e_{\min}$ for the first canonical vector and $e_{\max}$ for the last canonical vector. For any i we write

$$e_{\min+i} = \text{Next}^i \cdot e_{\min}$$

$$e_{\max+i} = \text{Prev}^i \cdot e_{\max}$$

We use the extra $\{0, 1\}^n \rightarrow \{0, 1\}^n$ functions that have a `for-MATLANG` translation:

$$\min(e) = \begin{cases} 1 & \text{if } e = e_{\min} \\ 0 & \text{otherwise} \end{cases}$$

$$\max(e) = \begin{cases} 1 & \text{if } e = e_{\max} \\ 0 & \text{otherwise} \end{cases}$$

$$\text{succ}(b_i, b_j) = \begin{cases} 1 & \text{if } i \leq j \\ 0 & \text{otherwise} \end{cases}$$

When used in `for-MATLANG` these functions output `[0]` and `[1]`.

Now

$$e_V := e_{\max-2}$$

$$e_{G_{top}} := e_{\max-1}$$

$$e_{V_{top}} := e_{\max}$$

For a canonical vector, let

$$\text{Iden}(b_i) := \sum v. \text{succ}(v, b_i) \cdot (v \cdot v^T).$$

This matrix has ones in the diagonal up to position i marked by e_i . We define the following sub-matrices of X :

$$V_{top} := X \cdot e_{V_{top}}$$

$$V := \text{Iden}(V_{top}) \cdot X \cdot e_V$$

$$G_{top} := X \cdot e_{G_{top}}$$

$$G := \text{Iden}(G_{top}) \cdot X \cdot \text{Iden}(e_{\max-3})$$

For example, if we are in a step where $\mathcal{G}.\text{size} = \mathcal{V}.\text{size} + 1$ then

$$X = \begin{bmatrix} \text{id}_1 & v_1 & 0 & 0 \\ \text{id}_2 & v_2 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ \text{id}_{j-1} & v_{j-1} & 0 & 1 \\ \text{id}_j & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 \end{bmatrix}, G = \begin{bmatrix} \text{id}_1 & 0 & 0 & 0 \\ \text{id}_2 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ \text{id}_{j-1} & 0 & 0 & 0 \\ \text{id}_j & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 \end{bmatrix},$$

$$V = \begin{bmatrix} v_1 \\ v_2 \\ \vdots \\ v_{j-1} \\ 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, G_{top} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 0 \\ 1 \\ 0 \\ \vdots \\ 0 \end{bmatrix}, V_{top} = \begin{bmatrix} 0 \\ 0 \\ \vdots \\ 1 \\ 0 \\ 0 \\ \vdots \\ 0 \end{bmatrix}.$$

Here, V is a vector encoding the stack of values in X and G is a matrix encoding the stack of gates in X . Note that what is *over* the top of the stacks is always set to zero due to

$\text{Iden}(G_{top})$ and $\text{Iden}(V_{top})$. Also, note that G is of the same size as X . We sometimes omit the zeroes due to simplicity.

To set the initial state (algorithm 3 line 2) we define the `for-MATLANG` expression:

$$\text{START} := e_{\min} \cdot \text{getroot}()^T + e_{\min} \cdot e_{G_{top}}^T.$$

For the initialize step, we define the `for-MATLANG` expressions: `INIT_PLUS` (algorithm 1, lines 2, 3, 4), `INIT_PROD` (algorithm 1, lines 5, 6, 7), `CONST` (algorithm 1, lines 8, 9) and `INPUT` (algorithm 1, lines 10, 11):

$$\begin{aligned} \text{INIT_PLUS} &:= \text{isplus}(G^T \cdot G_{top}) \times \\ &\quad \left[G + (\text{Next} \cdot G_{top}) \cdot \text{getfirst}(G^T \cdot G_{top})^T + \right. \\ &\quad \left. \text{Next} \cdot G_{top} \cdot e_{G_{top}}^T + V \cdot e_V^T + \text{Next} \cdot V_{top} \cdot e_{V_{top}}^T \right] \\ \text{INIT_PROD} &:= \text{isprod}(G^T \cdot G_{top}) \times \\ &\quad \left[G + (\text{Next} \cdot G_{top}) \cdot \text{getfirst}(G^T \cdot G_{top})^T + \right. \\ &\quad \left. \text{Next} \cdot G_{top} \cdot e_{G_{top}}^T + (V + \text{Next} \cdot v_{top}) \cdot e_V^T + \text{Next} \cdot V_{top} \cdot e_{V_{top}}^T \right] \\ \text{CONST} &:= \text{isone}(G^T \cdot G_{top}) \times \left[G + (V + \text{Next} \cdot V_{top}) \cdot e_V^T + \text{Next} \cdot V_{top} \cdot e_{V_{top}}^T \right] \\ \text{INPUT} &:= \text{isinput}(G^T \cdot G_{top}) \times \\ &\quad \left[G + \left(V + \left(v^T \cdot \text{Next} \cdot V_{top} \cdot \text{getinput}(G^T \cdot G_{top})^T \right) \right) \cdot e_V^T + \right. \\ &\quad \left. \text{Next} \cdot V_{top} \cdot e_{V_{top}}^T \right] \end{aligned}$$

Where v is the matrix variable stated in the theorem, the one associated with the input A of the circuit. Here, $G^T \cdot G_{top}$ is to get the current id in the top of the stack. In `INIT_PLUS` we get the current stack G , we add $\text{Next} \cdot G_{top} \cdot \text{getfirst}(G^T \cdot G_{top})^T$ which is an $n \times n$ matrix with the first child of $G^T \cdot G_{top}$ in the next row. Then $\text{Next} \cdot G_{top} \cdot e_{G_{top}}^T$ adds $\text{Next} \cdot G_{top}$ to the $n - 1$ column to mark the gate we added as the top. Next, we do the

same with the values by adding $V \cdot e_V + \text{Next} \cdot V_{top} \cdot e_{V_{top}}^T$. The other expressions work in a similar way.

The `for-MATLANG` expression equivalent to algorithm 1 is

$$\text{INIT} := \text{INIT_PLUS} + \text{INIT_PROD} + \text{CONST} + \text{INPUT}.$$

The idea is to return the matrix for the next iteration. Recall that here $\mathcal{G}.\text{size} = \mathcal{V}.\text{size} + 1$. So, when the operation is `INPUT` or `CONST`, if we start with

$$\begin{bmatrix} \text{id}_1 & v_1 & 0 & 0 \\ \text{id}_2 & v_2 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ \text{id}_{j-1} & v_{j-1} & 0 & 1 \\ \text{id}_j & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 \end{bmatrix}, \text{ then we return } \begin{bmatrix} \text{id}_1 & v_1 & 0 & 0 \\ \text{id}_2 & v_2 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ \text{id}_{j-1} & v_{j-1} & 0 & 0 \\ \text{id}_j & v_j & 1 & 1 \\ 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

When the operation is `INIT_PLUS` or `INIT_PROD`, if we start with

$$\begin{bmatrix} \text{id}_1 & v_1 & 0 & 0 \\ \text{id}_2 & v_2 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ \text{id}_{j-1} & v_{j-1} & 0 & 1 \\ \text{id}_j & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 \end{bmatrix}, \text{ then we return } \begin{bmatrix} \text{id}_1 & v_1 & 0 & 0 \\ \text{id}_2 & v_2 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ \text{id}_{j-1} & v_{j-1} & 0 & 0 \\ \text{id}_j & v_j & 0 & 1 \\ \text{id}_{j+1} & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 \\ \vdots & \vdots & \vdots & \vdots \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

For the aggregate expression (algorithm 2) we do the following. Let

$$\mathbb{E}[b_i, c] = \Sigma v.(v^T \cdot b_i) \cdot c \cdot v \cdot v^T + (1 - v^T \cdot b_i) \cdot v \cdot v^T,$$

namely, it is the identity with c in position (i, i) .

We define the expressions: AGG_PLUS (algorithm 2, lines 4, 5), AGG_PROD (algorithm 2, lines 6, 7), IS_NOT_LAST (algorithm 2, lines 8, 9), IS_LAST and POP:

$$\begin{aligned} \text{POP} &:= \text{Iden}(\text{Prev} \cdot G_{top}) \cdot G + \text{Prev} \cdot V_{top} \cdot e_{V_{top}}^T \\ \text{AGG_PLUS} &:= \text{isplus}(G^T \cdot (P \cdot G_{top})) \times \\ &\quad \left[(\text{Iden}(\text{Prev} \cdot V_{top}) \cdot V + (V^T \cdot V_{top}) (\text{Prev} \cdot V_{top})) \cdot e_V^T \right] \\ \text{AGG_PROD} &:= \text{isprod}(G^T \cdot (P \cdot G_{top})) \times \\ &\quad \left[(\mathbb{E}[\text{Prev} \cdot V_{top}, V^T \cdot V_{top}] \cdot \text{Iden}(\text{Prev} \cdot V_{top}) \cdot V) \cdot e_V^T \right] \\ \text{IS_NOT_LAST} &:= \text{not_last}(G^T \cdot (P \cdot G_{top}), G^T \cdot G_{top}) \times \\ &\quad \left[G_{top} \cdot \text{next_gate}(G^T \cdot (\text{Prev} \cdot G_{top}), G^T \cdot G_{top})^T + G_{top} \cdot e_{G_{top}}^T \right] \\ \text{IS_LAST} &:= (1 - \text{not_last}(G^T \cdot (P \cdot G_{top}), G^T \cdot G_{top})) \times \left[(\text{Prev} \cdot G_{top}) \cdot e_{G_{top}}^T \right] \end{aligned}$$

The `for-MATLANG` expression equivalent to algorithm 2 is

$$\text{AGG} := \text{POP} + \text{AGG_PLUS} + \text{AGG_PROD} + \text{IS_NOT_LAST} + \text{IS_LAST}.$$

The *Evaluate* method (algorithm 3) is defined as follows:

$$\begin{aligned}
\text{EVAL}[v] := & e_{\min}^T \cdot \{ \text{for } X, v_1, \dots, v_k \cdot : \\
& (\Pi_{i=1}^k \min(v_i)) \times \text{START} + \\
& (1 - \Pi_{i=1}^k \min(v_i)) \times \left((1 - \min(G_{\text{top}}) \cdot \min(V_{\text{top}})) \times \right. \\
& \quad \left. [(1 - G_{\text{top}}^T \cdot V_{\text{top}}) \times \text{INIT} + (G_{\text{top}}^T \cdot V_{\text{top}}) \times \text{AGG}] + \right. \\
& \quad \left. \min(G_{\text{top}}) \times \min(V_{\text{top}}) \times X \right) \\
& \} \cdot e_V
\end{aligned}$$

Note that the `for`-expression does the evaluation. The final output is in $X[1, \text{max}-2]$, we extract this value by multiplying the final result as $e_{\min}^T \cdot [\text{for}(\dots)] \cdot e_V$. We use $\Pi_{j=1}^k \min(v_i)$ to initialize with START .

Finally, we need to take care of all $n < n_0$, where $(\star)$ does not necessarily hold. For any i , let:

$\text{Eval}[i, A] :=$ the 1×1 matrix with the value of the polynomial $\Phi_n(A)$ when $n = i$.

Then we define:

$$\begin{aligned}
\Phi_n(a_1, \dots, a_i) := & \sum_{i=0}^{n_0-1} (e_{\min+i}^T \cdot (e_{\text{diag}}(e_1(v)) \cdot e_{\max})) \times \text{EVAL}[i, v] + \\
& ((\text{Next}^{n_0} \cdot e_{\min})^T \cdot e_1(v)) \times \text{EVAL}[v].
\end{aligned}$$

Above, $(e_{\min+i}^T \cdot e_{\max})$ checks if the dimension is equal to i (we multiply by the $n \times n$ identity $e_{\text{diag}}(e_1(v))$ to ensure typing), and $(\text{Next}^{n_0} \cdot e_{\min})^T \cdot e_1(e_{\min})$ checks if the dimension is greater or equal than n_0 . $\square$

While Theorem 5.1 gives us an idea on how to simulate arithmetic circuits, it does not tell us which classes of functions over real numbers can be computed by `for-MATLANG` expressions. In order to answer this question, we note that arithmetic circuits can be used to compute functions over real numbers. Formally, a circuit family $\{\Phi_n \mid n = 1, 2, \dots\}$ computes a function $f : \bigcup_{n \geq 1} \mathbb{R}^n \mapsto \mathbb{R}$, if for any $a_1, \dots, a_n \in \mathbb{R}$ it holds that

$\Phi_n(a_1, \dots, a_n) = f(a_1, \dots, a_n)$. To make the connection with `for-MATLANG`, we need to look at circuit families of bounded degree.

A circuit family $\{\Phi_n \mid n = 1, 2, \dots\}$ is said to be of *polynomial degree* if $\text{degree}(\Phi_n) \in O(p(n))$, for some polynomial $p(n)$. Note that polynomial size circuit families are not necessarily of polynomial degree. An easy corollary of Theorem 5.1 tells us that all functions computed by uniform family of circuits of polynomial degree and logarithmic depth can be simulated using `for-MATLANG` expressions. However, we can actually drop the restriction on circuit depth due to the result of Valiant et. al. (Valiant & Skyum, 1981) and Allender et. al. (Allender, Jiao, Mahajan, & Vinay, 1998) which says that any function computed by a uniform circuit family of polynomial degree (and polynomial depth), can also be computed by a uniform circuit family of logarithmic depth. Using this fact, we can conclude the following:

Corollary 5.1. *For any function f computed by a uniform family of arithmetic circuits of polynomial degree, there is an equivalent `for-MATLANG` formula e_f .*

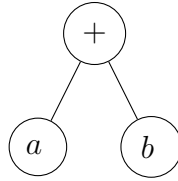
Note that there is nothing special about circuits that have a single output, and both Theorem 5.1 and Corollary 5.1 also hold for functions $f : \bigcup_{n \geq 1} \mathbb{R}^n \mapsto \mathbb{R}^{s(n)}$, where s is a polynomial. Namely, in this case, we can assume that circuits for f have multiple output gates, and that the depth reduction procedure of (Allender et al., 1998) is carried out for each output gate separately. Similarly, the construction underlying the proof of Theorem 5.1 can be performed for each output gate independently, and later composed into a single output vector.

5.3. From `for-MATLANG` to circuits

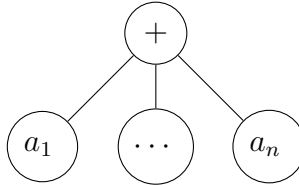
Now that we know that arithmetic circuits can be simulated using `for-MATLANG` expressions, it is natural to ask whether the same holds in the other direction. That is, we are asking whether for each `for-MATLANG` expression e over some schema $\mathcal{S}$ there is a

uniform family of arithmetic circuits computing precisely the same result depending on the input size.

In order to handle the fact that `for-MATLANG` expressions can produce any matrix, and not just a single value, as their output, we need to consider circuits which have multiple output gates. Similarly, we need to encode matrix inputs of a `for-MATLANG` expression in our circuits. We will write $\Phi(A_1, \dots, A_k)$, where Φ is an arithmetic circuit with multiple output gates, and each A_i is a matrix of dimensions $\alpha_i \times \beta_i$, with $\alpha_i, \beta_i \in \{n, 1\}$ to denote the input matrices for a circuit Φ . We will also write $\text{type}(\Phi) = (\alpha, \beta)$, with $\alpha, \beta \in \{n, 1\}$, to denote the size of the output matrix for Φ . We call such circuits *arithmetic circuits over matrices*. When $\{\Phi_n \mid n = 1, 2, \dots\}$ is a uniform family of arithmetic circuits over matrices, we will assume that the Turing machine for generating Φ_n also gives us the information about how to access a position of each input matrix, and how to access the positions of the output matrix, as is usually done when handling matrices with arithmetic circuits (Raz, 2003). The former will be denoted as $\Phi_n[i, j]$ when $\text{type}_S(\Phi) = (n, n)$, $\Phi_n[i, 1]$ when $\text{type}_S(\Phi) = (n, 1)$, $\Phi_n[1, j]$ when $\text{type}_S(\Phi) = (1, n)$ and Φ_n when $\text{type}_S(\Phi) = (1, 1)$. The notion of degree is extended to be the sum of the degrees of all the output gates. Also, when we write $a \oplus b$ we mean



When we write $\bigoplus_{l=1}^n a_l$ we mean



Same with $\otimes$. With this at hand, we can now show the following result.

Theorem 5.2. *Let e be a `for-MATLANG` expression over a schema $\mathcal{S}$, and let $V_1, \dots, V_k$ be the variables of e such that $\text{type}_{\mathcal{S}}(V_i) \in \{(\alpha, \alpha), (\alpha, 1), (1, \alpha), (1, 1)\}$. Then there exists a uniform arithmetic circuit family over matrices $\Phi_n(A_1, \dots, A_k)$ such that:*

- *For any instance $\mathcal{I} = (\mathcal{D}, \text{mat})$ such that $\mathcal{D}(\alpha) = n$ and $\text{mat}(V_i) = A_i$ it holds that:*
- $\llbracket e \rrbracket(\mathcal{I}) = \Phi_n(A_1, \dots, A_k)$.

PROOF. Let e be a `for-MATLANG` expression.

If $e = V$ then $\Phi_n^e := \Phi(A)$, and we have that

- If $\text{type}_{\mathcal{S}}(V) = (1, 1)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (1, 1)$ and Φ_n^e has the one input/output gate.
- If $\text{type}_{\mathcal{S}}(V) = (1, \alpha)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (1, n)$ and Φ_n^e has n input/output gates.
- If $\text{type}_{\mathcal{S}}(V) = (\alpha, 1)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (n, 1)$ and Φ_n^e has n input/output gates.
- If $\text{type}_{\mathcal{S}}(V) = (\alpha, \alpha)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (n, n)$ and Φ_n^e has n^2 input/output gates.

If $e = e'^T$ then $\Phi_n^e = \Phi_n^{e'}$.

- If $\text{type}_{\mathcal{S}}(\Phi_n^{e'}) = (1, 1)$ then $\Phi_n^e = \Phi_n^{e'}$ and $\text{type}_{\mathcal{S}}(\Phi_n^e) = (1, 1)$.
- If $\text{type}_{\mathcal{S}}(\Phi_n^{e'}) = (1, n)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (n, 1)$ and $\Phi_n^e[i, 1] := \Phi_n^{e'}[1, i]$.
- If $\text{type}_{\mathcal{S}}(\Phi_n^{e'}) = (n, 1)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (1, n)$ and $\Phi_n^e[1, j] := \Phi_n^{e'}[j, 1]$.
- If $\text{type}_{\mathcal{S}}(\Phi_n^{e'}) = (n, n)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (n, n)$ and $\Phi_n^e[i, j] := \Phi_n^{e'}[j, i]$.

If $e = \mathbf{1}(e')$ where $\text{type}_{\mathcal{S}}(\Phi_n^{e'}) = (\alpha, \beta)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (\alpha, 1)$ and $\Phi_n^e[i, 1] := 1$.

If $e = e_1 + e_2$ we have

- When $\text{type}_S(\Phi_n^{e_1}) = \text{type}_S(\Phi_n^{e_2}) = (1, 1)$ then $\text{type}_S(\Phi_n^e) = (1, 1)$ and $\Phi_n^e := \Phi_n^{e_1} \oplus \Phi_n^{e_2}$.
- When $\text{type}_S(\Phi_n^{e_1}) = \text{type}_S(\Phi_n^{e_2}) = (1, n)$ then $\text{type}_S(\Phi_n^e) = (1, n)$ and $\Phi_n^e[1, j] := \Phi_n^{e_1}[1, j] \oplus \Phi_n^{e_2}[1, j]$.
- When $\text{type}_S(\Phi_n^{e_1}) = \text{type}_S(\Phi_n^{e_2}) = (n, 1)$ then $\text{type}_S(\Phi_n^e) = (n, 1)$ and $\Phi_n^e[i, 1] := \Phi_n^{e_1}[i, 1] \oplus \Phi_n^{e_2}[i, 1]$.
- When $\text{type}_S(\Phi_n^{e_1}) = \text{type}_S(\Phi_n^{e_2}) = (n, n)$ then $\text{type}_S(\Phi_n^e) = (n, n)$ and $\Phi_n^e[i, j] := \Phi_n^{e_1}[i, j] \oplus \Phi_n^{e_2}[i, j]$.

If $e = f(e_1, \dots, e_k)$ we have two cases

- When f is the function $f_\odot$ (recall that this function is definable in $\text{MATLANG}[\emptyset]$ by Lemma 3.1) then

- If $\text{type}_S(\Phi_n^{e_1}) = \dots = \text{type}_S(\Phi_n^{e_k}) = (1, 1)$ then $\Phi_n^e := \bigotimes_{l=1}^k \Phi_n^{e_l}$.
- If $\text{type}_S(\Phi_n^{e_1}) = \dots = \text{type}_S(\Phi_n^{e_k}) = (1, n)$ then

$$\Phi_n^e[1, j] := \bigotimes_{l=1}^k \Phi_n^{e_l}[1, j]$$

.

- If $\text{type}_S(\Phi_n^{e_1}) = \dots = \text{type}_S(\Phi_n^{e_k}) = (n, 1)$ then

$$\Phi_n^e[i, 1] := \bigotimes_{l=1}^k \Phi_n^{e_l}[i, 1]$$

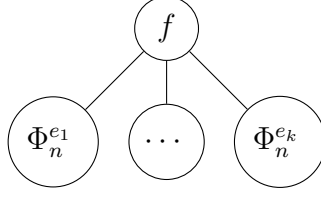
.

- If $\text{type}_S(\Phi_n^{e_1}) = \dots = \text{type}_S(\Phi_n^{e_k}) = (n, n)$ then

$$\Phi_n^e[i, j] := \bigotimes_{l=1}^k \Phi_n^{e_l}[i, j]$$

.

- When f is any other function, we prove the case when $\text{type}_S(\Phi_n^{e_1}) = \dots = \text{type}_S(\Phi_n^{e_k}) = (1, 1)$ (only case necessary, as discussed in Chapter 3.3.2). Here Φ_n^e is



Note that since for the context of this result we only consider $\text{for-MATLANG} = \text{MATLANG}[\emptyset]$, this case is not strictly necessary, modulo for $f_{\odot}, f_{\oplus}$ due to Lemma 3.1. However, if we extend the circuits with the same functions allowed in for-MATLANG , then our inductive construction still goes through, as just illustrated.

If $e = e_1 \cdot e_2$ we have

- When $\text{type}_{\mathcal{S}}(\Phi_n^{e_1}) = (1, 1)$ and $\text{type}_{\mathcal{S}}(\Phi_n^{e_2}) = (1, 1)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (1, 1)$ and $\Phi_n^e := \Phi_n^{e_1} \otimes \Phi_n^{e_2}$.
- When $\text{type}_{\mathcal{S}}(\Phi_n^{e_1}) = (1, 1)$ and $\text{type}_{\mathcal{S}}(\Phi_n^{e_2}) = (1, n)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (1, n)$ and $\Phi_n^e[1, j] := \Phi_n^{e_1} \otimes \Phi_n^{e_2}[1, j]$.
- When $\text{type}_{\mathcal{S}}(\Phi_n^{e_1}) = (n, 1)$ and $\text{type}_{\mathcal{S}}(\Phi_n^{e_2}) = (1, 1)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (n, 1)$ and $\Phi_n^e[i, 1] := \Phi_n^{e_1}[i, 1] \otimes \Phi_n^{e_2}$.
- When $\text{type}_{\mathcal{S}}(\Phi_n^{e_1}) = (n, 1)$ and $\text{type}_{\mathcal{S}}(\Phi_n^{e_2}) = (1, n)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (n, n)$ and $\Phi_n^e[i, j] := \Phi_n^{e_1}[i, 1] \otimes \Phi_n^{e_2}[1, j]$.
- When $\text{type}_{\mathcal{S}}(\Phi_n^{e_1}) = (1, n)$ and $\text{type}_{\mathcal{S}}(\Phi_n^{e_2}) = (n, 1)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (1, 1)$ and

$$\Phi_n^e := \bigoplus_{k=1}^n (\Phi_n^{e_1}[1, k] \otimes \Phi_n^{e_2}[k, 1]).$$

- When $\text{type}_{\mathcal{S}}(\Phi_n^{e_1}) = (1, n)$ and $\text{type}_{\mathcal{S}}(\Phi_n^{e_2}) = (n, n)$ then $\text{type}_{\mathcal{S}}(\Phi_n^e) = (1, n)$ and

$$\Phi_n^e[1, j] := \bigoplus_{k=1}^n (\Phi_n^{e_1}[1, k] \otimes \Phi_n^{e_2}[k, j]).$$

- When $\text{type}_S(\Phi_n^{e_1}) = (n, n)$ and $\text{type}_S(\Phi_n^{e_2}) = (n, 1)$ then $\text{type}_S(\Phi_n^e) = (n, 1)$ and

$$\Phi_n^e[i, 1] := \bigoplus_{k=1}^n (\Phi_n^{e_1}[i, k] \otimes \Phi_n^{e_2}[k, 1]).$$

- When $\text{type}_S(\Phi_n^{e_1}) = (n, n)$ and $\text{type}_S(\Phi_n^{e_2}) = (n, n)$ then $\text{type}_S(\Phi_n^e) = (n, n)$ and

$$\Phi_n^e[i, j] := \bigoplus_{k=1}^n (\Phi_n^{e_1}[i, k] \otimes \Phi_n^{e_2}[k, j]).$$

If $e = \text{for } X, v. e'(X, v)$, then define Φ^0 as the zero matrix circuit $\text{type}_S(\Phi^0) = (1, 1)$ if $\text{type}_S(\Phi_n^{e'}) = (1, 1)$ and $\text{type}_S(\Phi^0) = (n, n)$ if $\text{type}_S(\Phi_n^{e'}) = (n, n)$. Also, $\Phi^0 = 0$ and $\Phi^0[i, j] = 0 \forall i, j$ for each case respectively. Now for $i = 1, \dots, n$, define Φ^{v_i} as the circuit such that $\text{type}_S(\Phi^{v_i}) = (n, 1)$ and $\Phi^{v_i}[i, 1] := 1$ and zero otherwise. Finally, define

$$\Phi_n^e = \Phi_n^{e'} \left(\Phi_n^{e'} \left(\dots \left(\Phi_n^{e'} (\Phi^0, \Phi^{v_1}), \Phi^{v_2} \right) \dots, \Phi^{v_{n-1}} \right), \Phi^{v_n} \right).$$

Note that every circuit adds a constant number of layers except when in the case $e = \text{for } X, v. e'(X, v)$. This means that the depth still is polynomial. When $e = \text{for } X, v. e'(X, v)$ we have that the depth of the circuit is $n \cdot p(n)$, where the depth of $e'(X, v)$ is $p(n)$, so it also remains polynomial.

Here, we do not need to translate scalar multiplication because it can be simulated using the ones operator and $f_\odot$ (see Chapter 3.3.2).

Finally, we remark that when composing the circuits the fact that uniformity is preserved (i.e. the resulting circuit can be generated by a LOGSPACE machine) is proved analogously as when composing two LOGSPACE transducers (Arora & Barak, 2009). The only more involved case is treating for-loop construction, however, notice here that

we only need to keep track of where we are in the evaluation (i.e. which v_i we are processing), and not of all the previous results, given that they update the resulting matrix in a fixed order.

□

It is not difficult to see that the proof of Theorem 5.1 can also be extended to support arithmetic circuits over matrices. In order to identify the class of functions computed by `for-MATLANG` expressions, we need to impose one final restriction: than on the degree of an expression. Formally, the *degree of for-MATLANG expression e* over a schema $\mathcal{S}$, is the minimum of the degrees of any circuit family $\{\Phi_n \mid n = 1, 2, \dots\}$ that is equivalent to e . That is, the expression e is of polynomial degree, whenever there is an equivalent circuit family for e of a polynomial degree. For example, all `for-MATLANG` expressions seen so far have polynomial degree. With this definition, we can now identify the class of functions for which arithmetic circuits and `for-MATLANG` formulas are equivalent. This is the main technical contribution of the thesis.

Corollary 5.2. *Let f be a function with input matrices $A_1, \dots, A_k$ of dimensions $\alpha \times \beta$, with $\alpha, \beta \in \{n, 1\}$. Then, f is computed by a uniform circuit family over matrices of polynomial degree if and only if there is a `for-MATLANG` expression of polynomial degree for f .*

Note that this result crucially depends on the fact that expressions in `for-MATLANG` are of polynomial degree. Some `for-MATLANG` expressions are easily seen to produce results which are not polynomial. An example of such an expression is, for instance, $e_{\text{exp}} = \text{for } v, X = A. X \cdot X$, over a schema $\mathcal{S}$ with $\text{type}_{\mathcal{S}}(v) = (\gamma, 1)$, and $\text{type}_{\mathcal{S}}(X) = (1, 1)$. Over an instance which assigns n to γ this expression computes the function a^{2^n} , for $A = [a]$. Therefore, a natural question to ask then is whether we can determine the degree of a `for-MATLANG` expression. Unfortunately, as we show in the following proposition this question is in fact undecidable.

PROPOSITION 5.2. *Given a `for-MATLANG` expression e over a schema $\mathcal{S}$, it is undecidable to check whether e is of polynomial degree.*

PROOF. We show undecidability based on the following undecidable language:

$$\{\langle M \rangle \mid M \text{ is a deterministic TM which halts on the empty input}\},$$

where $\langle M \rangle$ is some string encoding of M . Consider a TM M described by $(Q, \Gamma = \{0, 1\}, q_0, q_m, \Delta)$ with $Q = \{q_1, \dots, q_m\}$ its states, q_1 being the initial state and q_m being the halting state, Γ is the tape alphabet, and Δ is a transition function from $Q \times \Gamma \rightarrow Q \times \Gamma \times \{\leftarrow, \sqcup, \rightarrow\}$. The simulation of linear space TM, as given in the proof of Proposition 5.1 can be easily modified to any TM M provided that we limit the execution of M to exactly n steps. Let e_M denote this expression. Similarly as in the linear space TM simulation, we have vector variables $Q_1, \dots, Q_m$ encoding the states, a single relation T encoding the tape and relation H_T encoding the position of the tape. When an instance $\mathcal{I}$ assigns n to α , we have a tape of length n at our disposal. This suffices if we let M run for n steps. We further observe that all vector variables can be assumed to be zero, initially. This is because we do not have input. So, let $\mathcal{I}_n^0$ denote the instance which assigns vector variables to the n -dimensional zero vector. Furthermore, by contrast to the linear space TM simulation, we use a single vector v (instead of k such vectors) to simulate n steps of M . Finally, we modify the expression given in the proof of Proposition 5.1 such $\llbracket e_M \rrbracket(\mathcal{I}_n^0)$ returns 1 if M halts in at most n steps, and 0 if M did not halt yet after n steps.

As a consequence, when M does not halt, $\llbracket e_M \rrbracket(\mathcal{I}_n^0) = 0$ for all $n \geq 0$. When M halts, there will be an n such that $\llbracket e_M \rrbracket(\mathcal{I}_n^0) = 1$. It now suffices to consider the `for-MATLANG` expression

$$d_M := e_M \cdot e_{\text{exp}}$$

where $e_{\text{exp}} = \text{for } v, X = \mathbf{1}(X)^T \cdot \mathbf{1}(X) \cdot X \cdot X$ such that $e_{\text{exp}}(\mathcal{I}_n^0) = n^{2^n}$. Then, when M does not halt we can clearly compute d_M with a constant degree circuit “0” for any n , otherwise, the circuit needed will be of exponential degree for at least one n , simply because no polynomial degree uniform circuit family can compute n^{2^n} . In other

words, deciding whether d_M has polynomial degree coincides with deciding whether M halts. □

Of course, one might wonder whether it is possible to define a syntactic subclass of `for-MATLANG` expressions that are of polynomial degree and can still express many important linear algebra algorithms. We identify one such class in Chapter 6.1, called `sum-MATLANG`, and in fact show that this class is powerful enough to capture relational algebra on (binary) K -relations.

5.4. Supporting additional operators

The equivalence of `for-MATLANG` and arithmetic circuits we prove above assumes that circuits can only use the sum and product gates (note that even without the sum and the product function, `for-MATLANG` can simulate these operations via matrix sum/product). However, both arithmetic circuits and expressions in `for-MATLANG` can be allowed to use a multitude of functions over $\mathbb{R}$. The most natural addition to the set of functions is the division operator, which is crucially needed in many linear algebra algorithms, such as, for instance, Gaussian elimination, or LU decomposition (recall Proposition 4.1). Interestingly, the equivalence in this case still holds, mainly due to a surprising result which shows that (almost all) divisions can in fact be removed from arithmetic circuits allowing sum, product, and division gates (Allender, 2004).

More precisely, in the work by Strassen, 1973; Borodin, von zur Gathen, & Hopcroft, 1982; Kaltofen, 1988 it was shown that for any function of the form $f = g/h$, where g and h are relatively prime polynomials of degree d , if f is computed by an arithmetic circuit of size s , then both g and h can be computed by a circuit whose size is polynomial in $s + d$. Given that we can postpone the division without affecting the final result, this, in essence, tells us that division can be eliminated (pushed to the top of the circuit), and we can work with sum-product circuits instead. The degree of a circuit for f , can then be defined as the maximum of degrees of circuits for g and h . Given this fact, we can again use the

depth reduction procedure of Allender et al., 1998, and extend Corollary 5.2 to circuits with division.

Corollary 5.3. *Let f be a function taking as its input matrices $A_1, \dots, A_k$ of dimensions $\alpha \times \beta$, with $\alpha, \beta \in \{n, 1\}$. Then, f is computed by a uniform circuit family over matrices of polynomial degree that allows divisions, if and only if there is a $\text{for-MATLANG}[f]$ expression of polynomial degree for f .*

An interesting line of future work here is to see which additional functions can be added to arithmetic circuits and for-MATLANG formulas, in order to preserve their equivalence. Note that this will crucially depend on the fact that these functions have to allow the depth reduction of Allender et al., 1998 in order to be supported.

6. RESTRICTING THE POWER OF FOR LOOPS

We conclude by zooming in on some special fragments of `for-MATLANG` and in which matrices can take values from an arbitrary (commutative) semiring K . In particular, we first consider `sum-MATLANG`, in which iterations can only perform additive updates, and show that it is equivalent in expressive power to the (positive) relational algebra on K -relations. We then extend `sum-MATLANG` such that also updates involving pointwise-multiplication (Hadamard product) are allowed. The resulting fragment, `FO-MATLANG`, is shown to be equivalent in expressive power to weighted logics. Finally, we consider the fragment `prod-MATLANG` in which updates involving sum and matrix multiplication, and possibly order information, is allowed. From the results in Chapter 4, we infer that the latter fragment suffices to compute matrix inversion. An overview of the fragments and their relationships are depicted in Figure 6.1.

6.1. Summation matlang and relational algebra

When defining 4-cliques and in several other expressions we have seen so far, we only update X by adding some matrix to it. This restricted form of for-loop proved useful throughout the thesis, and we therefore introduce it as a special operator. That is, we define:

$$\Sigma v.e := \text{for } v, X. X + e.$$

We define the subfragment of `for-MATLANG`, called `sum-MATLANG`, to consist of the Σ operator plus the “core” operators in `MATLANG`, namely, transposition, matrix multiplication and addition, scalar multiplication, and pointwise function applications.

One property of `sum-MATLANG` is that it only allows expressions of polynomial degree. Indeed, one can easily show that `sum-MATLANG` can only create matrix entries that are polynomial in the dimension n of the expression. More precisely, we can show the following:

PROPOSITION 6.1. *Every expression in sum-MATLANG is of polynomial degree.*

Interestingly enough, this restricted version of for-loop already allows us to capture the MATLANG operators that are not present in the syntax of sum-MATLANG . More precisely, we see from Examples 3.1 and 3.2 that the one-vector and diag operator are expressible in sum-MATLANG . Combined with the observation that the 4-clique expression of Example 3.3 is in sum-MATLANG , the following result is immediate.

Corollary 6.1. *MATLANG is strictly subsumed by sum-MATLANG .*

What operations over matrices can be defined with sum-MATLANG that is beyond MATLANG ? In the work of Brijder et al., 2020, it was shown that MATLANG is strictly included in the (positive) relational algebra on K -relations, denoted by RA_K^+ (Green, Karvounarakis, & Tannen, 2007).¹ It thus seems natural to compare the expressive power of sum-MATLANG with RA_K^+ . The main result in this chapter is that sum-MATLANG and RA_K^+ are equally expressive over binary schemas. To make this equivalence precise, we next give the definition of RA_K^+ (Green et al., 2007) and then show how to connect both formalisms.

Let $\mathbb{D}$ be a data domain and $\mathbb{A}$ a set of attributes. A relational signature is a finite subset of $\mathbb{A}$. A relational schema is a function $\mathcal{R}$ on finite set of symbols $\text{dom}(\mathcal{R})$ such that $\mathcal{R}(R)$ is a relation signature for each $R \in \text{dom}(\mathcal{R})$. To simplify the notation, from now on we write R to denote both the symbol R and the relational signature $\mathcal{R}(R)$. Furthermore, we write $R \in \mathcal{R}$ to say that R is a symbol of $\mathcal{R}$. For $R \in \mathcal{R}$, an R -tuple is a function $t : R \rightarrow \mathbb{D}$. We denote by $\text{tuples}(R)$ the set of all R -tuples. Given $X \subseteq R$, we denote by $t[X]$ the restriction of t to the set X .

A semiring $(K, \oplus, \odot, \mathbb{0}, 1)$ is an algebraic structure where K is a non-empty set, $\oplus$ and $\odot$ are binary operations over K , and $\mathbb{0}, 1 \in K$. Furthermore, $\oplus$ and $\odot$ are associative

¹The algebra used by Brijder et al., 2020 differs slightly from the one given by Green et al., 2007. In this thesis we work with the original algebra RA_K^+ as defined by Green et al., 2007.

operations, $\mathbb{0}$ and $\mathbb{1}$ are the identities of $\oplus$ and $\odot$ respectively, $\oplus$ is a commutative operation, $\odot$ distributes over $\oplus$, and $\mathbb{0}$ annihilates K (i.e. $\mathbb{0} \odot k = k \odot \mathbb{0} = \mathbb{0}$). As usual, we assume that all semirings in this thesis are commutative, namely, $\odot$ is also commutative. We use $\bigoplus_X$ or $\bigodot_X$ for the $\oplus$ - or $\odot$ -operation over all elements in X , respectively. Typical examples of semirings are the reals $(\mathbb{R}, +, \times, 0, 1)$, the natural numbers $(\mathbb{N}, +, \times, 0, 1)$, and the boolean semiring $(\{0, 1\}, \vee, \wedge, 0, 1)$.

Fix a semiring $(K, \oplus, \odot, \mathbb{0}, \mathbb{1})$ and a relational schema $\mathcal{R}$. A K -relation of $R \in \mathcal{R}$ is a function $r : \mathbf{tuples}(R) \rightarrow K$ such that the support $\text{supp}(r) = \{t \in \mathbf{tuples}(R) \mid r(t) \neq \mathbb{0}\}$ is finite. A K -instance $\mathcal{J}$ of $\mathcal{R}$ is a function that assigns relational signatures of $\mathcal{R}$ to K -relations. Given $R \in \mathcal{R}$, we denote by $R^{\mathcal{J}}$ the K -relation associated to R . Recall that $R^{\mathcal{J}}$ is a function and hence $R^{\mathcal{J}}(t)$ is the value in K assigned to t . Given a K -relation r we denote by $\mathbf{adom}(r)$ the active domain of r defined as $\mathbf{adom}(r) = \{t(a) \mid t \in \text{supp}(r) \wedge a \in R\}$. Then the active domain of an K -instance $\mathcal{J}$ of $\mathcal{R}$ is defined as $\mathbf{adom}(\mathcal{J}) = \bigcup_{R \in \mathcal{R}} \mathbf{adom}(R^{\mathcal{J}})$.

An RA_K^+ expression Q over $\mathcal{R}$ is given by the following syntax:

$$Q := R \mid Q \cup Q' \mid \pi_X(Q) \mid \sigma_X(Q) \mid \rho_f(Q) \mid Q \bowtie Q'$$

where $R \in \mathcal{R}$, $X \subseteq \mathbb{A}$ is finite, and $f : X \rightarrow Y$ is a one to one mapping with $Y \subseteq \mathbb{A}$. One can extend the schema $\mathcal{R}$ to any expression over $\mathcal{R}$ recursively as follows: $\mathcal{R}(R) = R$, $\mathcal{R}(Q \cup Q') = \mathcal{R}(Q) \cup \mathcal{R}(Q')$, $\mathcal{R}(\pi_X(Q)) = X$, $\mathcal{R}(\sigma_X(Q)) = \mathcal{R}(Q)$, $\mathcal{R}(\rho_f(Q)) = Y$ where $f : X \rightarrow Y$, and $\mathcal{R}(Q \bowtie Q') = \mathcal{R}(Q) \bowtie \mathcal{R}(Q')$ for every expressions Q and Q' . We further assume that any expression Q satisfies the following syntactic restrictions: $\mathcal{R}(Q') = \mathcal{R}(Q'')$ whenever $Q = Q' \cup Q''$, $X \subseteq \mathcal{R}(Q')$ whenever $Q = \pi_X(Q')$ or $Q = \sigma_X(Q')$, and $Y = \mathcal{R}(Q')$ whenever $Q = \rho_f(Q')$ with $f : X \rightarrow Y$.

Given an RA_K^+ expression Q and a K -instance $\mathcal{J}$ of $\mathcal{R}$, we define the semantics $\llbracket Q \rrbracket_{\mathcal{J}}$ as a K -relation of $\mathcal{R}(Q)$ as follows. For $X \subseteq \mathbb{A}$, let $\text{Eq}_X(t) = \mathbb{1}$ when $t(a) = t(b)$ for

every $a, b \in X$, and $\text{Eq}_X(t) = \mathbb{0}$ otherwise. For every tuple $t \in \mathcal{R}(Q)$:

$$\begin{aligned}
&\text{if } Q = R, \text{ then} && \llbracket Q \rrbracket_{\mathcal{J}}(t) = R^{\mathcal{J}}(t) \\
&\text{if } Q = Q_1 \cup Q_2, \text{ then} && \llbracket Q \rrbracket_{\mathcal{J}}(t) = \llbracket Q_1 \rrbracket_{\mathcal{J}}(t) \oplus \llbracket Q_2 \rrbracket_{\mathcal{J}}(t) \\
&\text{if } Q = \pi_X(Q'), \text{ then} && \llbracket Q \rrbracket_{\mathcal{J}}(t) = \bigoplus_{t': t'[X]=t} \llbracket Q' \rrbracket_{\mathcal{J}}(t') \\
&\text{if } Q = \sigma_X(Q'), \text{ then} && \llbracket Q \rrbracket_{\mathcal{J}}(t) = \llbracket Q' \rrbracket_{\mathcal{J}}(t) \odot \text{Eq}_X(t) \\
&\text{if } Q = \rho_f(Q'), \text{ then} && \llbracket Q \rrbracket_{\mathcal{J}}(t) = \llbracket Q' \rrbracket_{\mathcal{J}}(t \circ f) \\
&\text{if } Q = Q_1 \bowtie Q_2, \text{ then} && \llbracket Q \rrbracket_{\mathcal{J}}(t) = \llbracket Q_1 \rrbracket_{\mathcal{J}}(t[Y]) \odot \llbracket Q_2 \rrbracket_{\mathcal{J}}(t[Z]),
\end{aligned}$$

where $Y = \mathcal{R}(Q_1)$ and $Z = \mathcal{R}(Q_2)$. It is important to note that the $\bigoplus$ -operation in the semantics of $\pi_X(Q')$ is well-defined given that the support of $\llbracket Q' \rrbracket_{\mathcal{J}}$ is always finite.

We are now ready for comparing `sum-MATLANG` with RA_K^+ . First of all, we need to extend `sum-MATLANG` from $\mathbb{R}$ to any semiring. Let $\text{Mat}[K]$ denote the set of all K -matrices. Similarly as for `MATLANG` over $\mathbb{R}$, given a `MATLANG` schema $\mathcal{S}$, a K -instance $\mathcal{I}$ over $\mathcal{S}$ is a pair $\mathcal{I} = (\mathcal{D}, \text{mat})$, where $\mathcal{D} : \text{Symb} \mapsto \mathbb{N}$ assigns a value to each size symbol, and $\text{mat} : \mathcal{M} \mapsto \text{Mat}[K]$ assigns a concrete K -matrix to each matrix variable. Then it is straightforward to extend the semantics of `MATLANG`, `for-MATLANG`, and `sum-MATLANG` from $(\mathbb{R}, +, \times, 0, 1)$ to $(K, \oplus, \odot, \mathbb{0}, 1)$ by switching $+$ with $\oplus$ and $\times$ with $\odot$.

The next step to compare `sum-MATLANG` with RA_K^+ is to represent K -matrices as K -relations. Let $\mathcal{S} = (\mathcal{M}, \text{size})$ be a `MATLANG` schema. On the relational side we have for each size symbol $\alpha \in \text{Symb} \setminus \{1\}$, attributes α , row_α , and col_α in $\mathbb{A}$. Furthermore, for each $V \in \mathcal{M}$ and $\alpha \in \text{Symb}$ we denote by R_V and R_α its corresponding relation name, respectively. Then, given $\mathcal{S}$ we define the relational schema $\text{Rel}(\mathcal{S})$ such that $\text{dom}(\text{Rel}(\mathcal{S})) = \{R_\alpha \mid \alpha \in \text{Symb}\} \cup \{R_V \mid V \in \mathcal{M}\}$ where $\text{Rel}(\mathcal{S})(R_\alpha) = \{\alpha\}$ and:

$$\text{Rel}(\mathcal{S})(R_V) = \begin{cases} \{\text{row}_\alpha, \text{col}_\beta\} & \text{if } \text{size}(V) = (\alpha, \beta) \\ \{\text{row}_\alpha\} & \text{if } \text{size}(V) = (\alpha, 1) \\ \{\text{col}_\beta\} & \text{if } \text{size}(V) = (1, \beta) \\ \{\} & \text{if } \text{size}(V) = (1, 1). \end{cases}$$

Consider now a matrix instance $\mathcal{I} = (\mathcal{D}, \text{mat})$ over $\mathcal{S}$. Let $V \in \mathcal{M}$ with $\text{size}(V) = (\alpha, \beta)$ and let $\text{mat}(V)$ be its corresponding K -matrix of dimension $\mathcal{D}(\alpha) \times \mathcal{D}(\beta)$. To encode $\mathcal{I}$ as a K -instance in RA_K^+ , we use as data domain $\mathbb{D} = \mathbb{N} \setminus \{0\}$. Then we construct the K -instance $\text{Rel}(\mathcal{I})$ such that for each $V \in \mathcal{M}$ we define $R_V^{\text{Rel}(\mathcal{I})}(t) := \text{mat}(V)_{ij}$ whenever $t(\text{row}_\alpha) = i \leq \mathcal{D}(\alpha)$ and $t(\text{col}_\beta) = j \leq \mathcal{D}(\beta)$, and $\mathbb{0}$ otherwise. Furthermore, for each $\alpha \in \text{Symb}$ we define $R_\alpha^{\text{Rel}(\mathcal{I})}(t) := \mathbb{1}$ whenever $t(\alpha) \leq \mathcal{D}(\alpha)$, and $\mathbb{0}$ otherwise. In other words, R_α and R_β encodes the active domain of a matrix variable V with $\text{size}(V) = (\alpha, \beta)$. Given that the RA_K^+ framework of Green et al., 2007 represents the “absence” of a tuple in the relation with $\mathbb{0}$, we need to separately encode the indexes in a matrix. This is where $R_\alpha^{\text{Rel}(\mathcal{I})}$ and $R_\beta^{\text{Rel}(\mathcal{I})}$ are used for. We are now ready to state the first connection between sum-MATLANG and RA_K^+ by using the previous encoding. The proof of the proposition below is by induction on the structure of expressions.

PROPOSITION 6.2. *For each sum-MATLANG expression e over schema $\mathcal{S}$ such that $\mathcal{S}(e) = (\alpha, \beta)$ with $\alpha \neq 1 \neq \beta$, there exists an RA_K^+ expression $\Phi(e)$ over relational schema $\text{Rel}(\mathcal{S})$ such that $\text{Rel}(\mathcal{S})(\Phi(e)) = \{\text{row}_\alpha, \text{row}_\beta\}$ and such that for any instance $\mathcal{I}$ over $\mathcal{S}$,*

$$\llbracket e \rrbracket(\mathcal{I})_{i,j} = \llbracket \Phi(e) \rrbracket_{\text{Rel}(\mathcal{I})}(t)$$

for tuple $t(\text{row}_\alpha) = i$ and $t(\text{col}_\beta) = j$. Similarly for when e has schema $\mathcal{S}(e) = (\alpha, 1)$, $\mathcal{S}(e) = (1, \beta)$ or $\mathcal{S}(e) = (1, 1)$, then $\Phi(e)$ has schema $\text{Rel}(\mathcal{S})(\Phi(e)) = \{\text{row}_\alpha\}$, $\text{Rel}(\mathcal{S})(\Phi(e)) = \{\text{col}_\alpha\}$, or $\text{Rel}(\mathcal{S})(\Phi(e)) = \{\}$, respectively.

PROOF. We start from a matrix schema $\mathcal{S} = (\mathcal{M}, \text{size})$, where $\mathcal{M} \subset \mathcal{V}$ is a finite set of matrix variables, and $\text{size} : \mathcal{V} \mapsto \text{Symb} \times \text{Symb}$ is a function that maps each matrix variable to a pair of size symbols. On the relational side we have for each size symbol $\alpha \in \text{Symb} \setminus \{1\}$, attributes α , row_α , and col_α in $\mathbb{A}$. We also reserve some special attributes $\gamma_1, \gamma_2, \dots$ whose role will become clear shortly. For each $V \in \mathcal{M}$ and $\alpha \in \text{Symb}$ we denote by R_V and R_α its corresponding relation name, respectively.

Then, given $\mathcal{S}$ we define the relational schema $\text{Rel}(\mathcal{S})$ such that $\text{dom}(\text{Rel}(\mathcal{S})) = \{R_\alpha \mid \alpha \in \text{Symb}\} \cup \{R_V \mid V \in \mathcal{M}\}$ where $\text{Rel}(\mathcal{S})(R_\alpha) = \{\alpha\}$ and for all $V \in \mathcal{M}$:

$$\text{Rel}(\mathcal{S})(R_V) = \begin{cases} \{\text{row}_\alpha, \text{col}_\beta\} & \text{if } \text{size}(V) = (\alpha, \beta) \\ \{\text{row}_\alpha\} & \text{if } \text{size}(V) = (\alpha, 1) \\ \{\text{col}_\beta\} & \text{if } \text{size}(V) = (1, \beta) \\ \{\} & \text{if } \text{size}(V) = (1, 1). \end{cases}$$

Next, for a matrix instance $\mathcal{I} = (\mathcal{D}, \text{mat})$ over $\mathcal{S}$, let $V \in \mathcal{M}$ with $\text{size}(V) = (\alpha, \beta)$ and let $\text{mat}(V)$ be its corresponding K -matrix of dimension $\mathcal{D}(\alpha) \times \mathcal{D}(\beta)$. The K -instance in RA_K^+ according to $\mathcal{I}$ is $\text{Rel}(\mathcal{I})$ with data domain $\mathbb{D} = \mathbb{N} \setminus \{0\}$. For each $V \in \mathcal{M}$ we define $R_V^{\text{Rel}(\mathcal{I})}(t) := \text{mat}(V)_{ij}$ whenever $t(\text{row}_\alpha) = i \leq \mathcal{D}(\alpha)$ and $t(\text{col}_\beta) = j \leq \mathcal{D}(\beta)$, and 0 otherwise. Also, for each $\alpha \in \text{Symb}$ we define $R_\alpha^{\text{Rel}(\mathcal{I})}(t) := 1$ whenever $t(\alpha) \leq \mathcal{D}(\alpha)$, and 0 otherwise. If $\text{size}(V) = (\alpha, 1)$ then $R_V^{\text{Rel}(\mathcal{I})}(t) := \text{mat}(V)_{i1}$ whenever $t(\text{row}_\alpha) = i \leq \mathcal{D}(\alpha)$ and 0 otherwise. Similarly, if $\text{size}(V) = (1, \beta)$ then $R_V^{\text{Rel}(\mathcal{I})}(t) := \text{mat}(V)_{1j}$ whenever $t(\text{col}_\beta) = j \leq \mathcal{D}(\beta)$ and 0 otherwise. If $\text{size}(V) = (1, 1)$ then $R_V^{\text{Rel}(\mathcal{I})}(()) := \text{mat}(V)_{11}$.

Let e be a sum-MATLANG expression. In the following we need to distinguish between matrix variables v that occur in e as part of a sub-expression $\Sigma v.(\cdot)$, i.e., those variables that will be used to iterate over by means of canonical vectors, and those that are not. To make this distinction clear, we use $v_1, v_2, \dots$ for those “iterator” variables, and capital V for the other variables occurring in e . For simplicity, we assume that each occurrence of Σ has its own iterator variable associated with it.

We define free (iterator) variables, as follows. $\text{free}(V) := \emptyset$, $\text{free}(v) := \{v\}$, $\text{free}(e^T) := \text{free}(e)$, $\text{free}(e_1 + e_2) := \text{free}(e_1) \cup \text{free}(e_2)$, $\text{free}(e_1 \cdot e_2) := \text{free}(e_1) \cup \text{free}(e_2)$, $\text{free}(f_\odot(e_1, \dots, e_k)) := \text{free}(e_1) \cup \dots \cup \text{free}(e_k)$, and $\text{free}(e = \Sigma V.e_1) = \text{free}(e_1) \setminus \{v\}$. We will explicitly denote the free variables in an expression e by writing $e(v_1, \dots, v_k)$.

We now use the following induction hypotheses:

- If $e(v_1, \dots, v_k)$ is of type (α, β) then there exists a RA_K^+ expression Q such that $\text{Rel}(\mathcal{S})(Q(e)) = \{\text{row}_\alpha, \text{col}_\beta, \gamma_1, \dots, \gamma_k\}$ and such that

$$\llbracket Q(e) \rrbracket_{\text{Rel}(\mathcal{I})}(t) = \llbracket e \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_k \leftarrow b_{i_k}])_{i,j}$$

for tuple $t(\text{row}_\alpha) = i, t(\text{col}_\beta) = j$ and $t(\gamma_s) = i_s$ for $s = 1, \dots, k$.

- If $e(v_1, \dots, v_k)$ is of type $(\alpha, 1)$ then there exists a RA_K^+ expression Q such that $\text{Rel}(\mathcal{S})(Q(e)) = \{\text{row}_\alpha, \gamma_1, \dots, \gamma_k\}$ and such that

$$\llbracket Q(e) \rrbracket_{\text{Rel}(\mathcal{I})}(t) = \llbracket e \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_k \leftarrow b_{i_k}])_{i,1}$$

for tuple $t(\text{row}_\alpha) = i$, and $t(\gamma_s) = i_s$ for $s = 1, \dots, k$. And similarly for when e is type $(1, \beta)$.

- If $e(v_1, \dots, v_k)$ is of type $(1, 1)$ then there exists a RA_K^+ expression Q such that $\text{Rel}(\mathcal{S})(Q(e)) = \{\gamma_1, \dots, \gamma_k\}$ and such that

$$\llbracket Q(e) \rrbracket_{\text{Rel}(\mathcal{I})}(t) = \llbracket e \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_k \leftarrow b_{i_k}])_{1,1}$$

for tuple $t(\gamma_s) = i_s$ for $s = 1, \dots, k$.

Clearly, this suffices to show the proposition since we there consider expressions e for which $\text{free}(e) = \emptyset$, in which case the above statements reduce to the one given in the proposition.

The proof is by induction on the structure of `sum-MATLANG` expressions. In line with the simplifications in Chapter 3.3.2, it suffices to consider pointwise function application with $f_\odot$ instead of scalar multiplication. (We also note that we can express the one-vector operator in `sum-MATLANG`, so scalar multiplication can be expressed using $f_\odot$ in `sum-MATLANG`).

Let e be a `sum-MATLANG` expression.

- If $e = V$ then $Q(e) := R_V$.

- If $e = v_p$ then $Q(e) := \sigma_{\{\text{row}_\alpha, \gamma_p\}}(\rho_{\text{row}_\alpha \rightarrow \alpha}(R_\alpha) \bowtie \rho_{\gamma_p \rightarrow \alpha}(R_\alpha))$ when v_p is of type $(\alpha, 1)$. It is here that we introduce the attribute γ_p associated with iterator variable v_p . We note that

$$\llbracket Q(v_p) \rrbracket_{\text{Rel}(\mathcal{I})}(t) = \llbracket v_p \rrbracket(\mathcal{I}[v_p \leftarrow b_j])_{i,1} = (b_j)_{i,1}$$

for $t(\text{row}_\alpha) = i$ and $t[\gamma_p] = j$. Indeed, $(b_j)_{i,1} = \mathbb{1}$ if $j = i$ and this holds when $t(\text{row}_\alpha) = t[\gamma_p] = j$, and $(b_j)_{i,1} = \mathbb{0}$ if $j \neq i$ and this also holds when $t(\text{row}_\alpha) \neq t[\gamma_p] = j$.

- If $e(v_1, \dots, v_k) = (e_1(v_1, \dots, v_k))^T$ with $\mathcal{S}(e_1) = (\alpha, \beta)$ then

$$Q(e) := \begin{cases} \rho_{\text{row}_\alpha \rightarrow \text{col}_\alpha, \text{col}_\beta \rightarrow \text{row}_\beta}(Q(e_1)) & \text{if } \alpha \neq 1 \neq \beta; \\ \rho_{\text{row}_\alpha \rightarrow \text{col}_\alpha}(Q(e_1)) & \text{if } \alpha \neq 1 = \beta; \\ \rho_{\text{col}_\beta \rightarrow \text{row}_\beta}(Q(e_1)) & \text{if } \alpha = 1 \neq \beta; \\ Q(e_1) & \text{if } \alpha = 1 = \beta. \end{cases}$$

- If $e = e_1(v_1, \dots, v_k) + e_2(v_1, \dots, v_k)$ with $\mathcal{S}(e_1) = \mathcal{S}(e_2) = (\alpha, \beta)$ then $Q(e) := Q(e_1) \cup Q(e_2)$. We assume here that e_1 and e_2 have the same free variables. This is without loss of generality. Indeed, as an example, suppose that we have $e_1(v_1, v_2)$ and $e_2(v_2, v_3)$. Then, we can replace e_1 by $e_1(v_1, v_2, v_3) = (v_3^T \cdot v_3) \times e_1(v_1, v_2)$ and similarly, e_2 by $e_2(v_1, v_2, v_3) = (v_1^T \cdot v_1) \times e_2(v_2, v_3)$, where in addition we replace scalar multiplication with its simulation using $f_\odot$ and the ones vector, as mentioned earlier.
- If $e = f_\odot(e_1, \dots, e_k)$ with $\mathcal{S}(e_i) = \mathcal{S}(e_j)$ for all $i, j \in [1, k]$, then $Q(e) := Q(e_1) \bowtie \dots \bowtie Q(e_k)$.
- If $e = e_1 \cdot e_2$ with $\mathcal{S}(e_1) = (\alpha, \gamma)$ and $\mathcal{S}(e_2) = (\gamma, \beta)$, we have two cases. If $\gamma = 1$ then $Q(e) := Q(e_1) \bowtie Q(e_2)$. If $\gamma \neq 1$ then

$$Q(e) := \pi_{\{\text{row}_\alpha, \text{col}_\beta, \gamma_1, \dots, \gamma_k\}}(\rho_{C \rightarrow \text{col}_\gamma}(Q(e_1)) \bowtie \rho_{C \rightarrow \text{row}_\gamma}(Q(e_2))),$$

when

$$\text{Rel}(\mathcal{S})(Q(e_1)) = \{\text{row}_\alpha, \text{col}_\gamma, \gamma'_1, \dots, \gamma'_\ell\},$$

$$\text{Rel}(\mathcal{S})(Q(e_2)) = \{\text{row}_\gamma, \text{col}_\beta, \gamma''_1, \dots, \gamma''_\ell\} \text{ and}$$

$$\{\gamma_1, \dots, \gamma_k\} = \{\gamma'_1, \dots, \gamma'_k, \gamma''_1, \dots, \gamma''_m\}$$

- If $e(v_1, \dots, v_{p-1}, v_{p+1}, \dots, v_k) = \Sigma v_p.e_1(v_1, \dots, v_k)$ where $\mathcal{S}(e_1) = (\alpha, \beta)$ and $\mathcal{S}(V) = (\gamma, 1)$. Then we do

$$Q(e) := \pi_{\text{Rel}(\mathcal{S})(Q(e_1)) \setminus \{\gamma_p\}} Q(e_1).$$

Indeed, by induction we know that

$$\llbracket Q(e_1) \rrbracket_{\text{Rel}(\mathcal{I})}(t) = \llbracket e \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_k \leftarrow b_{i_k}])_{i,j}$$

for tuple $t(\text{row}_\alpha) = i, t(\text{col}_\beta) = j$ and $t(\gamma_s) = i_s$ for $s = 1, \dots, k$. Hence, for $t(\text{row}_\alpha) = i, t(\text{col}_\beta) = j$ and $t(\gamma_s) = i_s$ for $s = 1, \dots, k$ and $s \neq p$,

$$\llbracket Q(e_1) \rrbracket_{\text{Rel}(\mathcal{I})}(t) := \bigoplus_{i_p=1, \dots, \mathcal{D}(\gamma)} \llbracket e_1 \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_k \leftarrow b_{i_k}])_{i,j},$$

which precisely corresponds to

$$\llbracket \Sigma v_p.e_1(v_1, \dots, v_k) \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_{p-1} \leftarrow b_{p-1}, v_{p+1} \leftarrow b_{p+1}, \dots, v_k \leftarrow b_k])_{i,j}.$$

All other cases, when expressions have type $(\alpha, 1)$, $(1, \beta)$ or $(1, 1)$ can be dealt with in a similar way. $\square$

We now move to the other direction. To translate RA_K^+ into sum-MATLANG , we must restrict our comparison to RA_K^+ over K -relations with at most two attributes. Given that linear algebra works over vector and matrices, it is reasonable to restrict to unary or binary relations as input. Note that this is only a restriction on the input relations and not on intermediate relations, namely, expressions can create relation signatures of arbitrary size from the binary input relations. Thus, from now we say that a relational schema $\mathcal{R}$ is binary if $|R| \leq 2$ for every $R \in \mathcal{R}$. We also make the assumption that there is an (arbitrary) order, denoted by $<$, on the attributes in $\mathbb{A}$. This is to identify which attributes correspond

to rows and columns when moving to matrices. Then, given that relations will be either unary or binary and there is an order on the attributes, we write $t = (v)$ or $t = (u, v)$ to denote a tuple over a unary or binary relation R , respectively, where u and v is the value of the first and second attribute with respect to $<$.

Consider a binary relational schema $\mathcal{R}$. For each $R \in \mathcal{R}$ we associate a matrix variable V_R such that, if R is a binary relational signature, then V_R represents a (square) matrix, if R is unary, then V_R represents a vector and if $|R| = 0$ then V_R represents a constant. Formally, fix a symbol $\alpha \in \text{Symb} \setminus \{1\}$. Let $\text{Mat}(\mathcal{R})$ denote the MATLANG schema $(\mathcal{M}_{\mathcal{R}}, \text{size}_{\mathcal{R}})$ such that $\mathcal{M}_{\mathcal{R}} = \{V_R \mid R \in \mathcal{R}\}$ and $\text{size}_{\mathcal{R}}(V_R) = (\alpha, \alpha)$ whenever $|R| = 2$, $\text{size}_{\mathcal{R}}(V_R) = (\alpha, 1)$ whenever $|R| = 1$ and $\text{size}_{\mathcal{R}}(V_R) = (1, 1)$ whenever $|R| = 0$. Let $\mathcal{J}$ be the K -instance of $\mathcal{R}$ and suppose that $\text{adom}(\mathcal{J}) = \{d_1, \dots, d_n\}$ is the active domain (with arbitrary order) of $\mathcal{J}$. Define the matrix instance $\text{Mat}(\mathcal{J}) = (\mathcal{D}_{\mathcal{J}}, \text{mat}_{\mathcal{J}})$ such that $\mathcal{D}_{\mathcal{J}}(\alpha) = n$, $\text{mat}_{\mathcal{J}}(V_R)_{i,j} = R^{\mathcal{J}}((d_i, d_j))$ whenever $|R| = 2$, $\text{mat}_{\mathcal{J}}(V_R)_i = R^{\mathcal{J}}((d_i))$ whenever $|R| = 1$, and $\text{mat}_{\mathcal{J}}(V_R)_{1,1} = R^{\mathcal{J}}$ whenever $|R| = 0$. Note that, although each K -relation can have a different active domain, we encode them as square matrices by considering the active domain of the K -instance.

We next translate RA_K^+ expressions in to sum-MATLANG expressions over an extended schema. More specifically, for each attribute $A \in \mathbb{A}$ we define a vector variable v_A of type $(\alpha, 1)$. Then for each RA_K^+ expression Q with attributes $A_1, \dots, A_k$ we define a sum-MATLANG expression $e_Q(v_{A_1}, \dots, v_{A_k})$ of type $(1, 1)$ such that the following inductive hypothesis holds:

$$\llbracket e_Q \rrbracket(\text{Mat}(\mathcal{J})[v_{A_1} \leftarrow b_{i_1}, \dots, v_{A_k} \leftarrow b_{i_k}]) = \llbracket Q \rrbracket_{\mathcal{J}}(t) \quad (*)$$

where $t(A_s) = i_s$ for $s = 1, \dots, k$. The proof of this claim follows by induction on the structure of expressions:

- If $Q = R$, then $e_Q := v_{A_1}^T \cdot V_R \cdot v_{A_2}$ if $\mathcal{R}(R) = \{A_1, A_2\}$ with $A_1 < A_2$; $e_Q := V_R^T \cdot v_A$ if $\mathcal{R}(R) = \{A\}$; and $e_Q := V_R$ if $\mathcal{R}(R) = \{\}$.

- If $Q = Q_1 \cup Q_2$ then $e_Q := e_{Q_1} + e_{Q_2}$.
- If $Q = \pi_Y(Q_1)$ for $Y \subseteq \mathcal{R}(Q_1)$ and $\{B_1, \dots, B_l\} = \mathcal{R}(Q_1) \setminus Y$ then

$$e_Q := \Sigma v_{B_1} \cdot \Sigma v_{B_2} \cdot \dots \Sigma v_{B_l} \cdot e_{Q_1}$$

- If $Q = \sigma_Y(Q_1)$ with $Y \subseteq \mathcal{R}(Q_1)$ then

$$e_Q := e_{Q_1} \cdot \prod_{A, B \in Y} (v_A^T \cdot v_B).$$

Here Π is the matrix multiplication of expressions of type $(1, 1)$.

- If $Q = \rho_{X \mapsto Y}(Q_1)$ then

$$e_Q := e_{Q_1}[v_B \leftarrow v_A \mid A \in X, B \in Y, A \mapsto B].$$

In other words, we rename variable v_B with variable v_A in all the expression e_{Q_1} .

- If $Q = Q_1 \bowtie Q_2$ then $e_Q := e_{Q_1} \cdot e_{Q_2}$ where the product is over expression of type $(1, 1)$.

One can check, by induction over the construction, that the inductive hypothesis $(*)$ holds in each case. Now we can obtain the following result.

PROPOSITION 6.3. *Let $\mathcal{R}$ be a binary relational schema. For each RA_K^+ expression Q over $\mathcal{R}$ such that $|\mathcal{R}(Q)| = 2$, there exists a sum-MATLANG expression $\Psi(Q)$ over MATLANG schema $\text{Mat}(\mathcal{R})$ such that for any K -instance $\mathcal{J}$ with $\text{adom}(\mathcal{J}) = \{d_1, \dots, d_n\}$ over $\mathcal{R}$,*

$$\llbracket Q \rrbracket_{\mathcal{J}}((d_i, d_j)) = \llbracket \Psi(Q) \rrbracket(\text{Mat}(\mathcal{J}))_{i,j}.$$

Similarly for when $|\mathcal{R}(Q)| = 1$, or $|\mathcal{R}(Q)| = 0$ respectively.

PROOF. As a consequence of the previous discussion above, when Q is a RA_K^+ expression such that $\mathcal{R}(Q) = \{A_1, A_2\}$ with $A_1 < A_2$ then we define

$$\Psi(Q) = \Sigma v_{A_1} \cdot \Sigma v_{A_2} \cdot e_Q \cdot (v_{A_1} \cdot v_{A_2}^T).$$

Instead, when $\mathcal{R}(Q) = \{A\}$ we have

$$\Psi(Q) = \Sigma v_A \cdot (v_A \cdot e_Q).$$

And when $\mathcal{R}(Q) = \{\}$ we have

$$\Psi(Q) = e_Q.$$

By using the inductive hypothesis $(*)$ one can check that $\Psi(Q)$ works in each case as expected. $\square$

It is important to remark that the expression Q of the previous result can have intermediate expressions that are not necessary binary, given that the proposition only restricts that the input relation and the schema of Q must have arity at most two. We recall from the work of Brijder et al., 2020 that MATLANG corresponds to RA_K^+ where intermediate expressions are at most ternary, and this underlies, e.g., the inability of MATLANG to check for 4-cliques. In sum-MATLANG , we can deal with intermediate relations of arbitrary arity. In fact, each new attribute can be seen to correspond to an application of the Σ operator. For example, in the 4-clique expression, four Σ operators are needed, in analogy to how 4-clique is expressed in RA_K^+ .

Given the previous two propositions we derive the following conclusion which is the first characterization of relational algebra with a (sub)-fragment of linear algebra.

Corollary 6.2. *sum-MATLANG and RA_K^+ over binary relational schemas are equally expressive.*

As a direct consequence, we have that sum-MATLANG cannot compute matrix inversion. Indeed, using similar arguments as Brijder et al., 2019, i.e., by embedding RA_K^+ in

(infinitary) first-order logic with counting and by leveraging its locality, one can show that `sum-MATLANG` cannot compute the transitive closure of an adjacency matrix. By contrast, the transitive closure can be expressed by means of matrix inversion (Brijder et al., 2019). We also note that the evaluation of the Σ operator is independent of the order in which the canonical vectors are considered. This is because $\oplus$ is commutative. Hence, `sum-MATLANG` cannot express the order predicates mentioned in Chapter 3.

6.2. Hadamard product and weighted logics

Similarly to using `sum`, we can use other operations to update X in the `for`-loop. The next natural choice is to consider products of matrices. In contrast to matrix `sum`, we have two options: either we can choose to use matrix product or to use the pointwise matrix product, also called the Hadamard product. We treat matrix product in the next subchapter and first explain here the connection of `sum` and Hadamard product operators to weighted logics.

For the rest of this chapter, fix a semiring $(K, \oplus, \odot, \mathbb{0}, \mathbb{1})$. The Hadamard product over K -matrices can be defined as the pointwise application of $\odot$ between two matrices of the same size. Formally, we define the expression $e \circ e'$ where e, e' are expressions with respect to $\mathcal{S}$ and $\text{type}_{\mathcal{S}}(e) = \text{type}_{\mathcal{S}}(e')$ for some schema $\mathcal{S} = (\mathcal{M}, \text{size})$. Then the semantics of $e \circ e'$ is the pointwise application of $\odot$, namely, $\llbracket e \circ e' \rrbracket(\mathcal{I})_{ij} = \llbracket e \rrbracket(\mathcal{I})_{ij} \odot \llbracket e' \rrbracket(\mathcal{I})_{ij}$ for any instance $\mathcal{I}$ of $\mathcal{S}$. This enables us to define, similar as for Σv , the pointwise-product quantifier $\Pi^\circ v$ as follows:

$$\Pi^\circ v. e := \text{for } v, X = \mathbb{1}. X \circ e.$$

where $\mathbb{1}$ is a matrix with the same type as X and all entries equal to the $\mathbb{1}$ -element of K (i.e., we need to initialize X accordingly with the $\odot$ -operator). We call `FO-MATLANG` the subfragment of `for-MATLANG` that consists of `sum-MATLANG` extended with $\Pi^\circ v$.

Example 6.1. *Similar to the trace of a matrix, a useful function in linear algebra is to compute the product of the values on the diagonal. Using the $\Pi^\circ v$ operator, this can be easily expressed:*

$$e_{\text{dp}}(V) := \Pi^\circ v. v^T \cdot V \cdot v.$$

Clearly, the inclusion of this new operator extends the expressive power to `sum-MATLANG`. For example, $\llbracket e_{\text{dp}} \rrbracket(\mathcal{I})$ can be an exponentially large number in the dimension n of the input. By contrast, one can easily show that all expressions in `sum-MATLANG` can only return numbers polynomial in n . That is, `FO-MATLANG` is more expressive than `sum-MATLANG` and RA_K^+ .

To measure the expressive power of `FO-MATLANG`, we use weighted logics (Droste & Gastin, 2005) (WL) as a yardstick. Weighted logics extend monadic second-order logic from the boolean semiring to any semiring K . Furthermore, it has been used extensively to characterize the expressive power of weighted automata in terms of logic (Droste, Kuich, & Vogler, 2009). We use here the first-order subfragment of weighted logics to suit our purpose and, moreover, we extend its semantics over weighted structures (similar as Grädel & Tannen, 2017 do).

A relational vocabulary Γ is a finite collection of relation symbols such that each $R \in \Gamma$ has an associated arity, denoted by $\text{arity}(R)$. A K -weighted structure over Γ (or just structure) is a pair $\mathcal{A} = (A, \{R^{\mathcal{A}}\}_{R \in \Gamma})$ such that A is a non-empty finite set (i.e. the domain) and, for each $R \in \Gamma$, $R^{\mathcal{A}} : A^{\text{arity}(R)} \rightarrow K$ is a function that associates to each tuple in $A^{\text{arity}(R)}$ a weight in K .

Let X be a set of first-order variables. A K -weighted logic (WL) formula φ over Γ is defined by the following syntax:

$$\varphi := x = y \mid R(\bar{x}) \mid \varphi \oplus \varphi \mid \varphi \odot \varphi \mid \Sigma x. \varphi \mid \Pi x. \varphi$$

where $x, y \in X$, $R \in \Gamma$, and $\bar{x} = x_1, \dots, x_k$ is a sequence of variables in X such that $k = \text{arity}(R)$. As usual, we say that x is a free variable of φ , if x is not below Σx or Πx

quantifiers (e.g. x is free in $\Sigma y.R(x, y)$ but y is not). Given that K is fixed, from now on we talk about structures and formulas without mentioning K explicitly.

An assignment σ over a structure $\mathcal{A} = (A, \{R^A\}_{R \in \Gamma})$ is a function $\sigma : X \rightarrow A$. Given $x \in X$ and $a \in A$, we denote by $\sigma[x \mapsto a]$ a new assignment such that $\sigma[x \mapsto a](y) = a$ whenever $x = y$ and $\sigma[x \mapsto a](y) = \sigma(y)$ otherwise. For $\bar{x} = x_1, \dots, x_k$, we write $\sigma(\bar{x})$ to say $\sigma(x_1), \dots, \sigma(x_k)$. Given a structure $\mathcal{A} = (A, \{R^A\}_{R \in \Gamma})$ and an assignment σ , we define the semantics $\llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma)$ of φ as follows:

$$\begin{aligned} \text{if } \varphi &:= x = y, \text{ then} & \llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma) &= \begin{cases} 1 & \text{if } \sigma(x) = \sigma(y) \\ 0 & \text{otherwise} \end{cases} \\ \text{if } \varphi &:= R(\bar{x}), \text{ then} & \llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma) &= R^A(\sigma(\bar{x})) \\ \text{if } \varphi &:= \varphi_1 \oplus \varphi_2, \text{ then} & \llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma) &= \llbracket \varphi_1 \rrbracket_{\mathcal{A}}(\sigma) \oplus \llbracket \varphi_2 \rrbracket_{\mathcal{A}}(\sigma) \\ \text{if } \varphi &:= \varphi_1 \odot \varphi_2, \text{ then} & \llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma) &= \llbracket \varphi_1 \rrbracket_{\mathcal{A}}(\sigma) \odot \llbracket \varphi_2 \rrbracket_{\mathcal{A}}(\sigma) \\ \text{if } \varphi &:= \Sigma x. \varphi', \text{ then} & \llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma) &= \bigoplus_{a \in A} \llbracket \varphi' \rrbracket_{\mathcal{A}}(\sigma[x \mapsto a]) \\ \text{if } \varphi &:= \Pi x. \varphi', \text{ then} & \llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma) &= \bigodot_{a \in A} \llbracket \varphi' \rrbracket_{\mathcal{A}}(\sigma[x \mapsto a]) \end{aligned}$$

When φ contains no free variables, we omit σ and write $\llbracket \varphi \rrbracket_{\mathcal{A}}$ instead of $\llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma)$.

For comparing the expressive power of FO-MATLANG with WL, we have to show how to encode MATLANG instances into structures and vice versa. For this, we make two assumptions to put both languages at the same level: (1) we restrict structures to relation symbols of arity at most two and (2) we restrict instances to square matrices. The first assumption is for the same reasons as when comparing sum-MATLANG with RA_K^+ , and the second assumption is to have a crisp translation between both languages. Indeed, understanding the relation of FO-MATLANG with WL for non-square matrices is slightly more complicated and we leave this for future work.

Let $\mathcal{S} = (\mathcal{M}, \text{size})$ be a schema of square matrices, that is, there exists an α such that $\text{size}(V) \in \{1, \alpha\} \times \{1, \alpha\}$ for every $V \in \mathcal{M}$. Define the relational vocabulary $\text{WL}(\mathcal{S}) = \{R_V \mid V \in \mathcal{M}\}$ such that $\text{arity}(R_V) = 2$ if $\text{size}(V) = (\alpha, \alpha)$, $\text{arity}(R_V) = 1$ if $\text{size}(V) \in \{(\alpha, 1), (1, \alpha)\}$, and $\text{arity}(R_V) = 0$ otherwise. Then given a matrix

instance $\mathcal{I} = (\mathcal{D}, \text{mat})$ over $\mathcal{S}$ define the structure $\text{WL}(\mathcal{I}) = (\{1, \dots, n\}, \{R_V^{\mathcal{I}}\})$ such that $\mathcal{D}(\alpha) = n$ and $R_V^{\mathcal{I}}(i, j) = \text{mat}(V)_{i,j}$ if $\text{size}(V) = (\alpha, \alpha)$, $R_V^{\mathcal{I}}(i) = \text{mat}(V)_i$ if $\text{size}(V) \in \{(\alpha, 1), (1, \alpha)\}$, and $R_V^{\mathcal{I}} = \text{mat}(V)$ if $\text{size}(V) = (1, 1)$.

To encode weighted structures into matrices and vectors, the story is similar as for RA_K^+ . Let Γ be a relational vocabulary where $\text{arity}(R) \leq 2$. Define $\text{Mat}(\Gamma) = (\mathcal{M}_\Gamma, \text{size}_\Gamma)$ such that $\mathcal{M}_\Gamma = \{V_R \mid R \in \Gamma\}$ and $\text{size}_\Gamma(V_R)$ is equal to (α, α) , $(\alpha, 1)$, or $(1, 1)$ if $\text{arity}(R) = 2$, $\text{arity}(R) = 1$, or $\text{arity}(R) = 0$, respectively, for some $\alpha \in \text{Symb}$. Similarly, let $\mathcal{A} = (A, \{R^{\mathcal{A}}\}_{R \in \Gamma})$ be a structure with $A = \{a_1, \dots, a_n\}$, ordered arbitrarily. Then we define the matrix instance $\text{Mat}(\mathcal{A}) = (\mathcal{D}, \text{mat})$ such that $\mathcal{D}(\alpha) = n$, $\text{mat}(V_R)_{i,j} = R^{\mathcal{A}}(a_i, a_j)$ if $\text{arity}(R) = 2$, $\text{mat}(V_R)_i = R^{\mathcal{A}}(a_i)$ if $\text{arity}(R) = 1$, and $\text{mat}(V_R) = R^{\mathcal{A}}$ otherwise.

Let $\mathcal{S}$ be a MATLANG schema of square matrices and Γ a relational vocabulary of relational symbols of arity at most 2. We can then show the equivalence of FO-MATLANG and WL as follows.

PROPOSITION 6.4. *Weighted logics over Γ and FO-MATLANG over $\mathcal{S}$ have the same expressive power. More specifically,*

- *for each FO-MATLANG expression e over $\mathcal{S}$ such that $\mathcal{S}(e) = (1, 1)$, there exists a WL -formula $\Phi(e)$ over $\text{WL}(\mathcal{S})$ such that for every instance $\mathcal{I}$ of $\mathcal{S}$,*

$$\llbracket e \rrbracket(\mathcal{I}) = \llbracket \Phi(e) \rrbracket_{\text{WL}(\mathcal{I})}.$$

- *for each WL -formula φ over Γ without free variables, there exists a FO-MATLANG expression $\Psi(\varphi)$ such that for any structure $\mathcal{A}$ over $\text{Mat}(\Gamma)$,*

$$\llbracket \varphi \rrbracket_{\mathcal{A}} = \llbracket \Psi(\varphi) \rrbracket(\text{Mat}(\mathcal{A})).$$

PROOF. Both directions are proved by induction on the structure of expressions.

(FO-MATLANG to WL) First, let $\mathcal{S} = (\mathcal{M}, \text{size})$ be a schema of square matrices, that is, there exists an α such that $\text{size}(V) \in \{1, \alpha\} \times \{1, \alpha\}$ for every $V \in \mathcal{M}$. Define the relational vocabulary $\text{WL}(\mathcal{S}) = \{R_V \mid V \in \mathcal{M}\}$ such that $\text{arity}(R_V) = 2$ if $\text{size}(V) = (\alpha, \alpha)$, $\text{arity}(R_V) = 1$ if $\text{size}(V) \in \{(\alpha, 1), (1, \alpha)\}$, and $\text{arity}(R_V) = 0$ otherwise. Then given a matrix instance $\mathcal{I} = (\mathcal{D}, \text{mat})$ over $\mathcal{S}$ with $\mathcal{D}(\alpha) = n$ define the structure $\text{WL}(\mathcal{I}) = (\{1, \dots, n\}, \{R_V^{\mathcal{I}}\})$ such that $R_V^{\mathcal{I}}(i, j) = \text{mat}(V)_{i,j}$ if $\text{size}(V) = (\alpha, \alpha)$, $R_V^{\mathcal{I}}(i) = \text{mat}(V)_i$ if $\text{size}(V) \in \{(\alpha, 1), (1, \alpha)\}$, and $R_V^{\mathcal{I}} = \text{mat}(V)$ if $\text{size}(V) = (1, 1)$.

Similar to the proof of Proposition 6.2, for each expression $e(v_1, \dots, v_k)$ of type (α, α) we must encode in WL the α and the vector variables $v_1, \dots, v_k$. For this, we use variables x_α^{row} , x_α^{col} , and x_{v_i} for each variable $v_1, \dots, v_k$. Then we use the following inductive hypothesis (similar to Proposition 6.2):

- If $e(v_1, \dots, v_k)$ is of type (α, α) then there exists a WL formula

$$\varphi_e(x_\alpha^{\text{row}}, x_\alpha^{\text{col}}, x_{v_1}, \dots, x_{v_k})$$

such that

$$\llbracket \varphi_e \rrbracket_{\text{WL}(\mathcal{I})}(\sigma) = \llbracket e \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_k \leftarrow b_{i_k}])_{i,j}$$

for assignment σ with $\sigma(x_\alpha^{\text{row}}) = i$, $\sigma(x_\alpha^{\text{col}}) = j$ and $\sigma(x_{v_s}) = i_s$ for $s = 1, \dots, k$.

- If $e(v_1, \dots, v_k)$ is of type $(\alpha, 1)$ then there exists a WL formula

$$\varphi_e(x_\alpha^{\text{row}}, x_{v_1}, \dots, x_{v_k})$$

such that

$$\llbracket \varphi_e \rrbracket_{\text{WL}(\mathcal{I})}(\sigma) = \llbracket e \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_k \leftarrow b_{i_k}])_i$$

for assignment σ with $\sigma(x_\alpha^{\text{row}}) = i$ and $\sigma(x_{v_s}) = i_s$ for $s = 1, \dots, k$. And similarly for when e is type $(1, \alpha)$.

- If $e(v_1, \dots, v_k)$ is of type $(1, 1)$ then there exists a WL formula $\varphi_e(x_{v_1}, \dots, x_{v_k})$ such that

$$\llbracket \varphi_e \rrbracket_{\text{WL}(\mathcal{I})}(\sigma) = \llbracket e \rrbracket(\mathcal{I}[v_1 \leftarrow b_{i_1}, \dots, v_k \leftarrow b_{i_k}])$$

for assignment σ with $\sigma(x_{v_s}) = i_s$ for $s = 1, \dots, k$.

If we prove the previous statement we are done, because the last bullet is what we want to show when e has no free vector variables. Then rest of the proof is to go by induction on the structure of FO-MATLANG expressions. For a WL-formula φ and FO-variables x, y , we will write $\varphi[x \mapsto y]$ the formula φ when x is replaced with y all over the formula (syntactically). Let e be a FO-MATLANG expression.

- If $e := V$ and $\mathcal{S}(e) = (\alpha, \alpha)$ then $\varphi_e := R_V(x_\alpha^{\text{row}}, x_\alpha^{\text{col}})$. Similarly, if $\mathcal{S}(e)$ is of type $(\alpha, 1)$, $(1, \alpha)$, or $(1, 1)$, then $\varphi_e := R_V(x_\alpha^{\text{row}})$, $\varphi_e := R_V(x_\alpha^{\text{col}})$, and $\varphi_e := R_V$, respectively.
- If $e := v$, for $v \in \{v_1, \dots, v_k\}$, and $\mathcal{S}(v) = (\alpha, 1)$ then $\varphi_e := x_\alpha^{\text{row}} = x_v$. Similarly, if $\mathcal{S}(v) = (1, \alpha)$ then $\varphi_e := x_\alpha^{\text{col}} = x_v$.
- if $e := e_1^T$ and $\mathcal{S}(e) = (\alpha, \alpha)$ then

$$\varphi_e := \varphi_{e_1}[x_\alpha^{\text{row}} \mapsto x_\alpha^{\text{col}}, x_\alpha^{\text{col}} \mapsto x_\alpha^{\text{row}}].$$

Similarly, if $\mathcal{S}(e)$ is equal to $(\alpha, 1)$ or $(1, \alpha)$ then $\varphi_e := \varphi_{e_1}[x_\alpha^{\text{row}} \mapsto x_\alpha^{\text{col}}]$ and $\varphi_e := \varphi_{e_1}[x_\alpha^{\text{col}} \mapsto x_\alpha^{\text{row}}]$, respectively.

- If $e = e_1 + e_2$ with $\mathcal{S}(e_1) = \mathcal{S}(e_2)$, then $\varphi_e := \varphi_{e_1} \oplus \varphi_{e_2}$.
- If $e = f_\odot(e_1, \dots, e_k)$ with $\mathcal{S}(e_i) = \mathcal{S}(e_j)$ for all $i, j \in [1, k]$, then $\varphi_e := \varphi_{e_1} \odot \varphi_{e_2} \cdots \odot \varphi_{e_k}$.
- If $e = e_1 \cdot e_2$ with $\mathcal{S}(e_1) = \mathcal{S}(e_2) = (\alpha, \alpha)$, then $\varphi_e := \Sigma y. \varphi_{e_1}[x_\alpha^{\text{col}} \mapsto y] \odot \varphi_{e_2}[x_\alpha^{\text{row}} \mapsto y]$ where y is a fresh variable not mentioned in φ_{e_1} or φ_{e_2} .

Instead, if $\mathcal{S}(e_1) = (\alpha', 1)$ and $\mathcal{S}(e_2) = (1, \alpha'')$ with $\alpha', \alpha'' \in \{\alpha, 1\}$, then $\varphi_e := \varphi_{e_1} \odot \varphi_{e_2}$.

- If $e = \Sigma v.e_1(v)$, then we define $\varphi_e := \Sigma x_v. \varphi_{e_1}(x_v)$.
- If $e = \Pi^\circ v.e_1(v)$, then $\varphi_e := \Pi x_v. \varphi_{e_1}(x_v)$.

From the construction it is now straightforward to check that the inductive hypothesis holds for all cases. To conclude this direction, we have to define $\Phi(e) := \varphi_e$ for every expression e and we are done.

(WL to FO-MATLANG) We now encode weighted structures into matrices and vectors. Let Γ be a relational vocabulary where $\text{arity}(R) \leq 2$. Define $\text{Mat}(\Gamma) = (\mathcal{M}_\Gamma, \text{size}_\Gamma)$ such that $\mathcal{M}_\Gamma = \{V_R \mid R \in \Gamma\}$ and $\text{size}_\Gamma(V_R)$ is equal to (α, α) , $(\alpha, 1)$, or $(1, 1)$ if $\text{arity}(R) = 2$, $\text{arity}(R) = 1$, or $\text{arity}(R) = 0$, respectively, for some $\alpha \in \text{Symb}$. Similarly, let $\mathcal{A} = (A, \{R^{\mathcal{A}}\}_{R \in \Gamma})$ be a structure with $A = \{a_1, \dots, a_n\}$, ordered arbitrarily. Then we define the matrix instance $\text{Mat}(\mathcal{A}) = (\mathcal{D}, \text{mat})$ such that $\mathcal{D}(\alpha) = n$, $\text{mat}(V_R)_{i,j} = R^{\mathcal{A}}(a_i, a_j)$ if $\text{arity}(R) = 2$, $\text{mat}(V_R)_{i,1} = R^{\mathcal{A}}(a_i)$ if $\text{arity}(R) = 1$, and $\text{mat}(V_R)_{1,1} = R^{\mathcal{A}}$ otherwise.

Similar to the above direction, we have to encode the FO variables of a formula φ with vector variables in the equivalent FO-MATLANG expression e_φ . For this, for each FO variable x we define a vector variable v_x of type $(\alpha, 1)$. Then for each formula $\varphi(x_1, \dots, x_k)$ we define an expression $e_\varphi(v_{x_1}, \dots, v_{x_k})$ of type $(1, 1)$ such that for every assignment σ of $x_1, \dots, x_k$ we have:

$$\llbracket e_\varphi \rrbracket(\text{Mat}(\mathcal{A})[v_{x_1} \leftarrow b_{i_1}, \dots, v_{x_k} \leftarrow b_{i_k}]) = \llbracket \varphi \rrbracket_{\mathcal{A}}(\sigma)$$

such that $\sigma(x_s) = i_s$ for every $s \leq k$. Note that when the formula has no free variables, the proof of the proposition is shown. Finally, we proceed by induction over the formula φ over Γ .

- If $\varphi := x = y$, then $e_\varphi := v_x^T \cdot v_y$.

- If $\varphi := R(x, y)$, then $e_\varphi := v_x^T \cdot V_R \cdot v_y$. Similarly, if $\varphi := R(x)$ or $\varphi := R$, then $e_\varphi := V_R^T \cdot v_x$ and $e_\varphi := V_R$, respectively.
- If $\varphi = \varphi_1 \oplus \varphi_2$, then $e_\varphi := e_{\varphi_1} + e_{\varphi_2}$.
- If $\varphi = \varphi_1 \odot \varphi_2$, then $e_\varphi := f_\odot(e_{\varphi_1}, e_{\varphi_2})$.
- If $\varphi = \Sigma x. \varphi_1$, then $e_\varphi := \Sigma v_x. e_{\varphi_1}$.
- If $\varphi = \Pi^\circ x. \varphi_1$, then $e_\varphi := \Pi v_x. e_{\varphi_1}$.

The inductive hypothesis can be proved following the above construction. To finish the proof, we define $\Psi(\varphi) := e_\varphi$ and the proposition is shown. □

6.3. Matrix multiplication as a quantifier

In a similar way, we can consider a fragment in which sum and the usual product of matrices can be used in for-loops. Formally, for an expression e we define the operator:

$$\Pi v. e = \text{for } v, X = I \cdot X \cdot e.$$

where I is the identity matrix. We call `prod-MATLANG` the subfragment of `for-MATLANG` that consists of `sum-MATLANG` extended with Πv . It is readily verified that $\Pi^\circ v$ can be expressed in terms of Πv . Furthermore, by contrast to the Hadamard product, matrix multiplication is a non-commutative operator. As a consequence, one can formulate expressions that are not invariant under the order in which the canonical vectors are processed.

PROPOSITION 6.5. *Every expression in FO-MATLANG can be defined in prod-MATLANG. Moreover, there exists an expression that uses the Πv quantifier that cannot be defined in FO-MATLANG.*

What is interesting is that `sum-MATLANG` extended with Πv suffices to compute the transitive closure, provided that we allow for the $f_{>0}$ function. Indeed, one can use the

expression $e_{\text{TC}}(V) := f_{>0}(\Pi v. (e_{\text{Id}} + V))$ for this purpose because $\llbracket e_{\text{TC}} \rrbracket(\mathcal{I}) = f_{>0}((I + A)^n)$ when $\mathcal{I}$ assigns an $n \times n$ adjacency matrix A to V , and non-zero entries in $(I + A)^n$ coincide with non-zero entries in the transitive closure of A . Furthermore, if we extend this fragment with access to the matrix $S_{<}$, defining the (strict) order on canonical vectors, then Csanky's matrix inversion algorithm becomes expressible (if $f_{/}$ is allowed). We leave the study of this fragment and, in particular, the relationship to full `for-MATLANG`, for future work.

We conclude by verifying that the fragment defined in Chapter 6.3, i.e, `prod-MATLANG` extended with order and $f_{>0}$, can perform matrix inversion and compute the determinant. To this aim, we verify that all order predicates in Chapter 3.2.2 can be derived using Σ , Π , $f_{>0}$ and $e_{S_{<}}$. Given this, it suffices to observe that Csanky's algorithm, as shown in Chapter 4.2, only relies on expressions using Σ and Π and order information on canonical vectors and $f_{/}$. As consequence, our fragment can perform matrix inversion and compute the determinant.

It remains to show that if we have $e_{S_{<}}$, using Σ and Π and $f_{>0}$ we can define all order predicates from Chapter 3.2.2. We note that due to the restricted for-loops in Σ and Π , we do not have access to the intermediate result in the iterations and as such, it is unclear whether order information can be computed. This is why we assume access to $e_{S_{<}}$.

We first remark that if we have $e_{S_{<}}$, we can also obtain $e_{S_{\leq}}$ by adding e_{Id} . Hence, we can compute `succ` and `succ+` as well. Furthermore,

$$e_{\min} := \Sigma v. [\Pi w. \text{succ}(w, v)] \times v.$$

$$e_{\max} := \Sigma v. [\Pi w. (1 - \text{succ}(w, v))] \times v.$$

Both expressions are only using Σ and Π and `succ`, so are in our fragment. Furthermore, if we have $f_{>0}$ then we can define

$$e_{\text{Pred}} := e_{S_{<}} - f_{>0}(e_{S_{<}}^2)$$

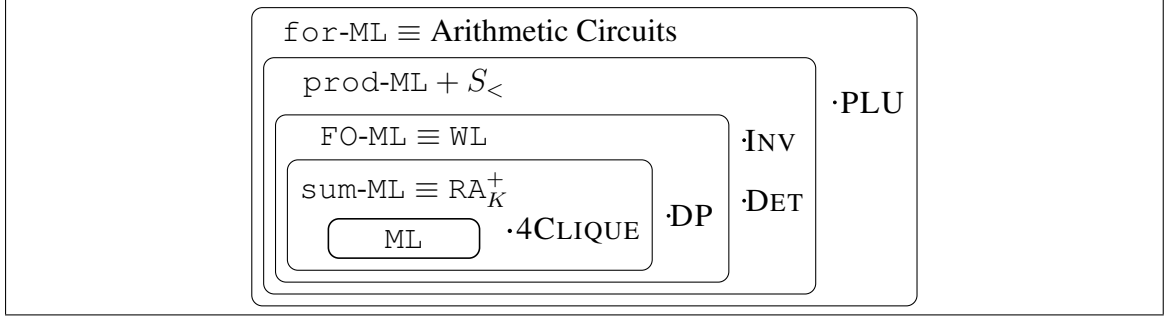


Figure 6.1. Fragments of `for-MATLANG` and their equivalences. The functions `4CLIQUE`, `DP` (diagonal product), `INV`, `DET`, and `PLU` decomposition are placed in their fragments.

Also, recall that $e_{\text{Next}} := e_{\text{Pred}}^T$. As a consequence, we can now define $\text{prev}(v)$ and $\text{next}(v)$ as in 3.2.2. Similarly, it is readily verified that also $e_{\text{getPrevMatrix}(V)}$, $e_{\text{getNextMatrix}(V)}$, $e_{\text{min}+i}$ and $e_{\text{max}+i}$ can be expressed in our fragment.

Finally, in Figure 6.1 we show a diagram of all the fragments of `for-MATLANG` introduced in this chapter and their corresponding equivalent formalisms.

7. CONCLUSIONS

The goal of this thesis was to create a matrix query language that has the expressive power to compute classical linear algebra algorithms on matrix data. Our focus was on language design and expressivity. A language simple in a linear algebra sense, but powerful enough to do complex matrix procedures.

Based on previous work, we proposed `for-MATLANG`, an extension of `MATLANG` with limited recursion, and showed that it is able to capture most of linear algebra due to its connection to arithmetic circuits.

Then, to further study the expressive power of `for-MATLANG`, we aim to restrict the introduced `for` operator. We focus on its data manipulation and the matrix operations allowed. This revealed interesting connections to relational algebra on annotated relations and weighted logics.

The former proves that `for-MATLANG` can describe how to perform useful operations on matrix data, which has many direct applications in machine learning algorithms, as well as graph operations on adjacency matrices.

Note that `for-MATLANG` works on any semiring, thus other logics or complexity classes may be captured by the language. We leave this open for future work.

Another interesting direction for future work relates to efficient evaluation of (fragments) of `for-MATLANG`. A possible starting point is the work of Christ, Demmel, Knight, Scanlon, & Yelick, 2013 in which a general methodology for communication-optimal algorithms for for-loop linear algebra programs is proposed. An efficient evaluation goes in hand with implementing data procedures in `for-MATLANG`.

Additionally, our results of Chapter 6.1 show interesting opportunities for exploring how relational joins can be evaluated using matrix operations, and vice versa, which can lead to an efficient implementation of relational algebra operations.

The query language `for-MATLANG` does not impose a way of actually computing the operations. For instance, the operators can be implemented as GPU procedures, known to accelerate computation on arrays of numbers. More complex algorithms on GPU architecture may require a generalization of `for-MATLANG` to tensor objects, an interesting future line of work.

This investigation aims to be a setting stone in the modern ways of how numerical data is treated, when it's stored as n -dimensional arrays. In the particular case of `for-MATLANG`, we treat two dimensional data at most. It defines core operations that have the expressiveness to do the main linear algebra operations. The goal in the implementation of `for-MATLANG` in an algorithm is to treat numerical data as a database itself, and let it run the needed operations on it's own.

The former could lead, on one hand, to efficient implementations in many fields of algorithm evaluation, or at least a more direct way of declaring the needed operation. On the other hand, heuristics can be applied in procedures that handle numerical data at some point. Maybe it could lead to a proper declarative language for GPU computing.

A more direct experimental way of future work is the implementation in recommendation systems, where matrix operations are present in most (if not all) algorithms. Here, user-item interaction data is handled as numerical values and the data is stored in matrices. The idea is that if we need a result of a procedure on a matrix, the algorithm asks for the result as a `for-MATLANG` expression.

REFERENCES

- Allender, E. (2004). Arithmetic circuits and counting complexity classes. *Complexity of Computations and Proofs, Quaderni di Matematica*, 13, 33–72.
- Allender, E., Jiao, J., Mahajan, M., & Vinay, V. (1998). Non-commutative arithmetic circuits: Depth reduction and size lower bounds. *Theor. Comput. Sci.*, 209(1-2), 47–86. Retrieved from [https://doi.org/10.1016/S0304-3975\(97\)00227-2](https://doi.org/10.1016/S0304-3975(97)00227-2)
- Arora, S., & Barak, B. (2009). *Complexity theory: A modern approach*. Cambridge University Press Cambridge.
- Barceló, P., Higuera, N., Pérez, J., & Subercaseaux, B. (2020). On the expressiveness of LARA: A unified language for linear and relational algebra. In *Proceedings of the 23rd international conference on database theory, ICDT* (pp. 6:1–6:20). Retrieved from <https://doi.org/10.4230/LIPIcs.ICDT.2020.6>
- Boehm, M., Kumar, A., & Yang, J. (2019). *Data management in machine learning systems*. Morgan & Claypool Publishers. Retrieved from <https://doi.org/10.2200/S00895ED1V01Y201901DTM057>
- Borodin, A., von zur Gathem, J., & Hopcroft, J. (1982). Fast parallel matrix and gcd computations. In *23rd annual symposium on foundations of computer science (sfcs 1982)* (pp. 65–71).
- Brijder, R., Geerts, F., den Bussche, J. V., & Weerwag, T. (2019). On the expressive power of query languages for matrices. *ACM Trans. Database Syst.*, 44(4), 15:1–15:31. Retrieved from <https://doi.org/10.1145/3331445>
- Brijder, R., Gyssens, M., & den Bussche, J. V. (2020). On matrices and k-relations. In *Foundations of information and knowledge systems (foiks)* (pp. 42–57).

Christ, M., Demmel, J., Knight, N., Scanlon, T., & Yelick, K. A. (2013). Communication lower bounds and optimal algorithms for programs that reference arrays - part 1. Retrieved from <http://dx.doi.org/10.21236/ada584726>

Cormen, T. H., Leiserson, C. E., Rivest, R. L., & Stein, C. (2001). *Introduction to algorithms* (2nd ed.). The MIT Press.

Csanky, L. (1976). Fast parallel matrix inversion algorithms. *SIAM J. Comput.*, 5(4), 618–623. Retrieved from <https://doi.org/10.1137/0205040>

Dawar, A., Grohe, M., Holm, B., & Laubner, B. (2009). Logics with rank operators. In *Proceedings of the 24th annual ieee symposium on logic in computer science* (pp. 113–122). (<https://doi.org/10.1109/LICS.2009.24>)

Droste, M., & Gastin, P. (2005). Weighted automata and weighted logics. In *Icalp* (Vol. 3580, pp. 513–525).

Droste, M., Kuich, W., & Vogler, H. (2009). *Handbook of weighted automata*. Springer Science & Business Media.

Geerts, F. (2019). On the expressive power of linear algebra on graphs. In *22nd international conference on database theory, ICDT 2019, march 26-28, 2019, lisbon, portugal* (pp. 7:1–7:19).

Geerts, F. (2020). When can matrix query languages discern matrices? In *23rd international conference on database theory, ICDT 2020, march 30-april 2, 2020, copenhagen, denmark* (Vol. 155, pp. 12:1–12:18). Retrieved from <https://doi.org/10.4230/LIPIcs.ICDT.2020.12>

Grädel, E., & Tannen, V. (2017). Semiring provenance for first-order model checking. *CoRR*, *abs/1712.01980*. Retrieved from <http://arxiv.org/abs/1712.01980>

Green, T. J., Karvounarakis, G., & Tannen, V. (2007). Provenance semirings. In *Pods* (pp.

31–40).

Grohe, M., & Pakusa, W. (2017). Descriptive complexity of linear equation systems and applications to propositional proof complexity. In *32nd annual ACM/IEEE symposium on logic in computer science* (pp. 1–12). (<https://doi.org/10.1109/LICS.2017.8005081>)

Hella, L., Libkin, L., Nurmonen, J., & Wong, L. (2001). Logics with aggregate operators. *J. ACM*, 48(4), 880–907. doi: 10.1145/502090.502100

Holm, B. (2010). *Descriptive complexity of linear algebra* (Unpublished doctoral dissertation). University of Cambridge.

Hutchison, D., Howe, B., & Suciu, D. (2017). Laradb: A minimalist kernel for linear and relational algebra computation. In *Proceedings of the 4th ACM SIGMOD workshop on algorithms and systems for mapreduce and beyond, beyondmr@sigmod 2017, chicago, il, usa, may 19, 2017* (pp. 2:1–2:10).

Jankov, D., Luo, S., Yuan, B., Cai, Z., Zou, J., Jermaine, C., & Gao, Z. J. (2019). Declarative recursive computation on an RDBMS. *Proc. VLDB Endow.*, 12(7), 822–835. Retrieved from <http://www.vldb.org/pvldb/vol12/p822-jankov.pdf>

Kaltofen, E. (1988). Greatest common divisors of polynomials given by straight-line programs. *Journal of the ACM (JACM)*, 35(1), 231–264.

Khamis, M. A., Ngo, H. Q., Nguyen, X., Olteanu, D., & Schleich, M. (2018). AC/DC: in-database learning thunderstruck. In S. Schelter, S. Seufert, & A. Kumar (Eds.), *Proceedings of the second workshop on data management for end-to-end machine learning* (pp. 8:1–8:10). ACM. Retrieved from <https://doi.org/10.1145/3209889.3209896>

Kunft, A., Alexandrov, A., Katsifodimos, A., & Markl, V. (2016). Bridging the gap: Towards optimization across linear and relational algebra. In *Proceedings of the 3rd acm*

sigmod workshop on algorithms and systems for mapreduce and beyond (pp. 1:1–1:4). (<http://doi.acm.org/10.1145/2926534.2926540>)

Luo, S., Gao, Z. J., Gubanov, M., Perez, L. L., & Jermaine, C. (2018). Scalable linear algebra on a relational database system. *SIGMOD Rec.*, 47(1), 24–31. (<http://doi.acm.org/10.1145/3277006.3277013>)

Press, W. H., Teukolsky, S. A., Vetterling, W. T., & Flannery, B. P. (1992). *Numerical recipes in c, 2nd edition*. Cambridge University Press. Retrieved from <http://numerical.recipes>

Raz, R. (2003). On the complexity of matrix product. *SIAM J. Comput.*, 32(5), 1356–1369. Retrieved from <https://doi.org/10.1137/S0097539702402147>

Shpilka, A., & Yehudayoff, A. (2010). Arithmetic circuits: A survey of recent results and open questions. *Foundations and Trends in Theoretical Computer Science*, 5(3-4), 207–388. Retrieved from <http://dx.doi.org/10.1561/04000000039>

Strassen, V. (1973). Vermeidung von divisionen. *Journal für die reine und angewandte Mathematik*, 264, 184–202.

Valiant, L. G., & Skyum, S. (1981). Fast parallel computation of polynomials using few processors. In *International symposium on mathematical foundations of computer science* (pp. 132–139).