

Practical and Flexible Equality Saturation

Max Willsey

A dissertation
submitted in partial fulfillment of the
requirements for the degree of

Doctor of Philosophy

University of Washington

2021

Reading Committee:

Luis Ceze, Chair

Adriana Schulz

Zachary Tatlock

Program Authorized to Offer Degree:
Computer Science & Engineering

© Copyright 2021

Max Willsey

University of Washington

Abstract

Practical and Flexible Equality Saturation

Max Willsey

Chair of the Supervisory Committee:

Luis Ceze

Paul G. Allen School of Computer Science & Engineering

Programming language tools like compilers, optimizers, verifiers, and synthesizers rely on term rewriting to effectively manipulate programs. While powerful and well-studied, term rewriting traditionally suffers from a critical stumbling block: users must *choose* when and how to apply the right rewrite, and the quality of the results hinges on this difficult decision. A recent technique called *equality saturation* mitigates this “rewrite choice” issue by allowing many rewrites to apply simultaneously. Despite its promise, the technique’s applicability has been limited by lack of flexibility and poor scalability. This thesis offers theoretical and practical advances that make equality saturation fast and flexible enough to use in real-world applications in any domain.

On the theoretical side, this work contributes two techniques to make e-graphs, the data structure underlying equality saturation, better suited to the algorithm’s needs. A new amortized invariant restoration technique called *rebuilding* takes advantage of equality saturation’s distinct workload, providing asymptotic speedups over current techniques in practice. A general mechanism called *e-class analyses* integrates domain-specific analyses into the e-graph, reducing the need for ad hoc manipulation.

We implemented these techniques in a new open-source library called egg. egg has been used

to achieve state-of-the-art results in many domains, including floating point accuracy, automatic vectorization, deep learning compute graphs, 3D CAD decompilation, and linear algebra kernels, among others. We present case studies that highlight how egg’s performance and flexibility helped these projects succeed, making the case that equality saturation is ready for a wide variety of real-world use cases.

Contents

1	Introduction	1
1.1	The World Before egg	2
1.2	The egg Era	4
2	Background	7
2.1	Term Representation and Rewriting	7
2.2	E-Graphs	10
2.2.1	Definitions	10
2.2.2	E-Graph Invariants	12
2.2.3	Interface and Rewriting	13
2.3	Equality Saturation	15
3	Rebuilding: A New Take on E-graph Invariant Maintenance	19
3.1	Upward Merging	20
3.2	Rebuilding in Detail	20
3.2.1	Examples of Rebuilding	23
3.2.2	Proof of Congruence	24
3.3	Rebuilding and Equality Saturation	25
3.4	Evaluating Rebuilding	27

4	Extending E-graphs with E-class Analyses	33
4.1	E-Class Analyses	34
4.1.1	Maintaining the Analysis Invariant	35
4.1.2	Example: Constant Folding	37
4.2	Conditional and Dynamic Rewrites	38
4.3	Extraction	39
5	egg: Easy, Extensible, and Efficient E-graphs	41
5.1	Ease of Use	42
5.2	Extensibility	44
5.3	Efficiency	47
6	Case Studies	49
6.1	Szalinski: Decompiling CAD into Structured Programs	50
6.1.1	Implementation	51
6.1.2	Inverse Transformations	53
6.1.3	Results	54
6.2	Herbie: Improving Floating Point Accuracy	55
6.2.1	Implementation	56
6.2.2	Results	57
6.3	Tensat: Optimizing Deep Learning Computation Graphs	59
6.3.1	Representation	60
6.3.2	Exploration Phase	61
6.3.3	Extraction Phase	62
6.3.4	Evaluation	66
6.4	Ruler: Rewrite Synthesis using Equality Saturation	71
6.4.1	Ruler’s Algorithm	73
6.4.2	Choosing Rules	75

6.4.3	Comparison with CVC4	78
6.4.4	Synthesizing Herbie Rewrites	79
7	Conclusion	83
7.1	Future Work	85
	Bibliography	87

Acknowledgments

This thesis covers work done in the final year and a half of my PhD. The work done before then, largely with the Molecular Information Systems Lab at UW, was and is a great source of pride. I am very fortunate to enjoy the endlessly supportive mentorship of my advisor Luis Ceze, who is the very first person you want on your team, and who introduced me to the collaborative, cross-disciplinary environment at MISL.

On the egg side of things, I owe the all the fun and any success I had to the entire PLSE group. Thanks especially to Zach Tatlock, without whom PLSE would simply not be the same. The main component of this thesis is the egg paper, and my co-authors in that work deserve a special shoutout for making that project what it is: Chandrakana Nandi, Remy Wang, Oliver Flatt, Zach Tatlock, and Pavel Panchekha.

Thanks to my friends and family for making this journey possible and even enjoyable. And finally, thank you Sami for doing this whole thing with me.

Chapter 1

Introduction

At the heart of programming languages lies the question of how to represent and manipulate programs. Nearly all aspects of programming language work—including theorem proving, optimizing compilers, and program synthesis—depend on these fundamental notions, since they dictate how (and how efficiently!) a tool stores and works with programs. The choice of how to represent and manipulate programs largely determines the efficacy of a tool or technique.

This thesis takes a fresh look at a data structure called the *e-graph* [Nel80] and a technique called *equality saturation* [TSTL09] for representing and manipulating programs, respectively. Together, they offer great promise over the status quo of syntax trees and term rewriting. Unfortunately, poor scalability and flexibility has limited the approach’s reach.

To bring e-graphs and equality saturation into the programming language practitioner’s toolbox, this thesis makes the following claim:

E-graphs and equality saturation are compelling techniques for program representation and manipulation that should now be considered for programming tools across many domains.

The rest of this document supports that claim up in three ways:

1. Our original research contributions presented in Chapter 3 and Chapter 4 advance the state-of-the-art around e-graphs and equality saturation, enhancing both so that together they form a compelling alternative to conventional term rewriting.
2. These theoretical advances are realized in `egg`, a first-of-its-kind tool that has brought e-graphs and equality saturation to a wider range of users and domains than ever before. Chapter 5 documents how `egg` combines novel techniques with myriad practical niceties to make a generic, high-performance implementation of e-graphs and equality saturation.
3. Chapter 6 makes an empirical case for the thesis statement in the form of published projects that rely on `egg`. These case studies demonstrate that e-graphs and equality saturation can now be used to achieve state-of-the-art results in various domains.

1.1 The World Before `egg`

Abstract syntax trees and directed term rewriting are (and will likely remain) the most popular approached for program representation and manipulation. We defer discussion of issues with this approach to Chapter 2. Here I would like to set the stage a little and describe the setting in which this thesis’s contributions were made.

Equality graphs (e-graphs) were developed in late 1970s to efficiently represent congruence relations in automated theorem provers (ATPs). At a high level, e-graphs [Nel80, NO05] extend union-find [Tar75] to compactly represent equivalence classes of expressions while maintaining a key invariant: the equivalence relation is closed under congruence.¹

¹Intuitively, congruence simply means that $a \equiv b$ implies $f(a) \equiv f(b)$.

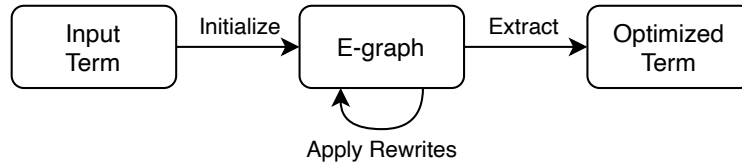


Figure 1.1: Equality saturation optimizes a program p by storing it in an e-graph, growing the e-graph into a large set of equivalent terms by applying rewrites, and finally selecting the best program that is equivalent to p .

In the 2000s, work like Denali [JNR02] and the first equality saturation papers [TSTL09, STL11] began to repurpose e-graphs as the basis for program optimization. Given an input program p , equality saturation constructs an e-graph E that represents a large set of programs equivalent to p , and then extracts the “best” program from E . The e-graph is grown by repeatedly applying pattern-based rewrites $\ell \rightarrow r$. Each rewrite $\ell \rightarrow r$ includes a pattern ℓ to match and a pattern r to instantiate and merge with the matched subterm. Critically, these rewrites only add information² to the e-graph, eliminating the need for careful ordering. Upon reaching a fixed point (*saturation*), E will represent *all equivalent ways* to express p with respect to the given rewrites. After saturation (or timeout), a final extraction procedure analyzes E and selects the optimal program according to a user-provided cost function.

Ideally, a user could simply provide a language grammar and rewrites, and equality saturation would produce a effective optimizer. Three challenges blocked this ideal:

1. Maintaining congruence can become expensive as E grows. In part, this is because e-graphs from the conventional ATP setting remained unspecialized to the distinct *equality saturation workload*. Early applications based on equality saturation had to limit their searches, which can impact the quality of results.
2. Many applications critically depend on *domain-specific analyses*, but integrating them

²As opposed to traditional term rewriting which only considers a single term at a time. Section 2.1 covers this in detail.

required ad hoc extensions to the e-graph. The lack of a general extension mechanism forced researchers to re-implement equality saturation from scratch several times [PSSWT15, TSTL09, WZN⁺19].

3. Even outside of the context of domain-specific analyses, the lack of a generic implementation of e-graphs and equality saturation made the “just bring your grammar and rewrites” vision impossible. Users had to start from scratch and reimplement state-of-the-art techniques for their own domain.

1.2 The egg Era

The work presented in this thesis addresses all of the concerns raised above. Our theoretical contributions make e-graphs and equality saturation faster and more flexible, and egg makes it all usable in real-world applications.

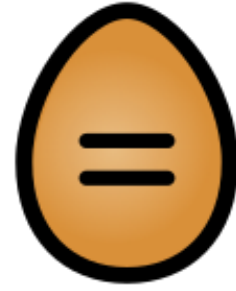


Figure 1.2: The egg logo.

Equality Saturation Workload ATPs frequently query and modify e-graphs and additionally require the ability to undo modifications (e.g., in DPLL(T) [DP60]). This forces conventional e-graph designs to maintain the congruence invariant after every operation. In contrast, the equality saturation workload can be factored into distinct phases of (1) querying the e-graph to simultaneously find all rewrite matches and (2) modifying the e-graph to merge in equivalences for all matched terms.

We present a new amortized algorithm called *rebuilding* (Chapter 3) that defers e-graph invariant maintenance to equality saturation phase boundaries without compromising soundness. Empirically, rebuilding provides asymptotic speedups over conventional approaches.

Domain-specific Analyses Equality saturation is primarily driven by syntactic rewriting, but many applications require additional interpreted reasoning to bring domain knowledge into the e-graph. Past implementations have resorted to ad hoc e-graph manipulations to integrate what would otherwise be simple program analyses like constant folding.

To flexibly incorporate such reasoning, we introduce a new, general mechanism called *e-class analyses* (Chapter 4). An e-class analysis annotates each e-class (an equivalence class of terms) with facts drawn from a semilattice domain. As the e-graph grows, facts are introduced, propagated, and joined to satisfy the *e-class analysis invariant*, which relates analysis facts to the terms represented in the e-graph. Rewrites cooperate with e-class analyses by depending on analysis facts and adding equivalences that in turn establish additional facts. The examples and case studies demonstrate e-class analyses like constant folding and free variable analysis which required bespoke customization in previous equality saturation implementations.

A generic, high-performance implementation We implement rebuilding and e-class analyses in an open-source³ library called **egg** (**e**-graphs **g**ood). **egg** specifically targets equality saturation, taking advantage of its workload characteristics and supporting easy extension mechanisms to provide e-graphs specialized for program synthesis and optimization. **egg** also addresses more prosaic challenges, e.g., parameterizing over user-defined languages, rewrites, and cost functions while still providing an optimized implementation. Our case studies in Chapter 6 demonstrate how **egg**’s features constitute a general, reusable e-graph library that can support equality saturation across diverse domains.

In summary, I believe that equality saturation has a big role in the future of programming languages. When so many of the challenges to building programming tools revolve around *choice*,

³ web: <https://egraphs-good.github.io>
source: <https://github.com/egraphs-good/egg>
documentation: <https://docs.rs/egg>

equality saturation's ability to operate over many terms *simultaneously* is hard to ignore. Even in situations where equality saturation is not the best approach, it may be the easiest, since it offloads work from the developer to the computer. More complex applications may require making a particular choice at some points and taking all possible paths in others; egg's practical, generic implementation is compatible with this approach as well.

Fast and flexible equality saturation provides the power to *not* choose how to manipulate programs, which can be a boon to both power users and domain experts who may not have expertise in building programming language tools. I hope this thesis and egg are steps toward having e-graphs and equality saturation in the toolbox of anyone working with programs, regardless of their expertise, the size and scale of the problem, or the application domain.

Chapter 2

Background

2.1 Term Representation and Rewriting

The most common and easily understood program representation is the venerable *abstract syntax tree* (AST). An AST is a recursive tree data structure where a node is an operator and some number—potentially zero—of children ASTs. For example, the program $(a * 2)/2$ is represented by the AST shown in Figure 2.1a. The operators $*$ and $/$ each take two children; the variable a and the number 2 are leaf operators that take no children.

ASTs are the primary data structures used in most programming tools, but they are not the only ones. A common alternate representation is the *term graph*, which can be viewed as a variant of ASTs that allows directed acyclic graphs instead of just trees. Term graphs can therefore capture sharing across common parts of a program (also known as a term). Figure 2.1b shows a term graph for the program $(a * 2)/2$.

Regardless of the choice of program representation, a common technique for program manipulation is rewriting. In this paradigm, program transformations are given as a set of rewrites, where each rewrite $\ell \rightarrow r$ specifies a pattern ℓ to search for and a another pattern r with which to



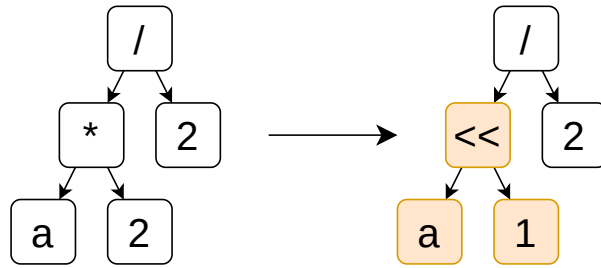
Figure 2.1: Different representations of the program $(a * 2)/2$ have different characteristics. The syntax tree is simpler (a), but the term graph is smaller since it captures sharing (b).

replace each instance of ℓ found in the program. For example, applying the rewrite $x * 2 \rightarrow x + x$ to our term $(a * 2)/2$ yields $(a + a)/2$.

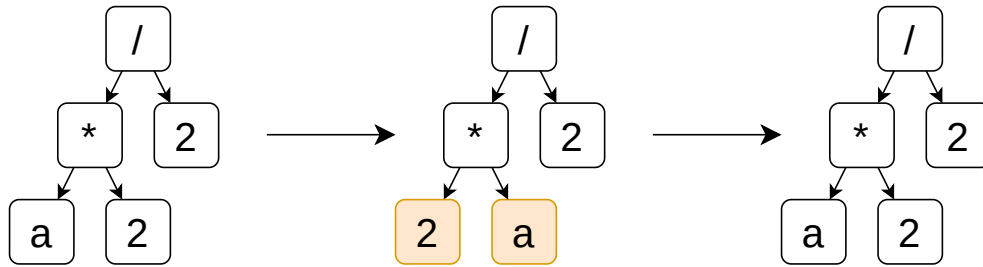
Stated more formally, applying rewrite $\ell \rightarrow r$ on a term t works as follows. First, search for the ℓ within t , yielding a substitution σ and a subterm s of t such that σ maps variables from ℓ to terms and $\ell[\sigma] = s$, where $\ell[\sigma]$ denotes replacing the variables in ℓ with the corresponding terms in σ . With the substitution in hand, applying the rewrite simply replaces $\ell[\sigma]$ with $r[\sigma]$ in t (since $\ell[\sigma]$ equals some subterm s of t).

Rewriting offers an intuitive, compositional, and efficient method to transform programs that is used in programming tools of all shapes and sizes. It is a well-researched technique with a wealth of literature (surveyed in [Der93, DJ90]), and many programming language tools implement and rely on it.

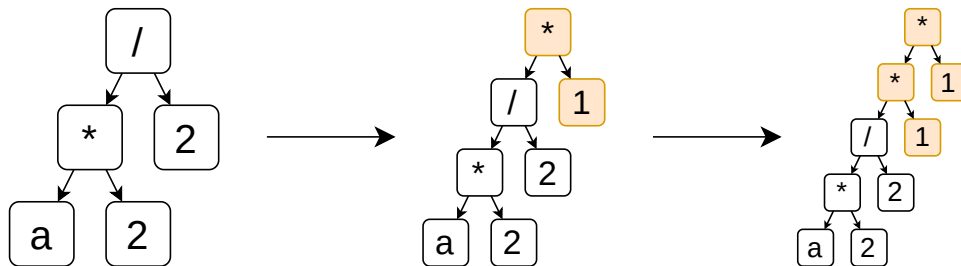
Term rewriting (in its traditional directed form) does, however, suffer from pitfalls that can complicate or prevent the building of certain rewrite-based systems. Many of these weaknesses boil down to the fact that term rewriting operates on *one term at a time*; once a term is rewritten, you are left with the new version and have essentially forgotten the old version. The quality of a rewriting system's output can heavily depend on the order and manner in which it applies its rewrites. In other words, *choices really matter* in this paradigm. The compilers community refers to this as the *phase-ordering* problem [WS97, TSTL09].



(a) Sometimes a locally “good” rewrite can be a poor choice down the road. In this case, we would ultimately like to rewrite $(a * 2)/2$ to 2 . But applying $x * 2 \rightarrow x \ll 1$ is hard to pass up, since replacing a (relatively) expensive multiplication with a cheap bitshift is nearly always a good decision. Applying that locally beneficial rewrite makes canceling out the 2s much more difficult.



(b) Seemingly innocuous rewrites like commutativity of multiplication ($x * y \rightarrow y * x$) can send a rewriting system into a loop. Directed rewriting systems can avoid this by applying these rewrites in only one direction or trying to observe when they have encountered a term they have seen before.



(c) Expansive rewrites like $x \rightarrow x * 1$ can enable other rewrites, but they are problematic for traditional rewriting systems since they can always be applied, potentially leading to infinitely large terms.

Figure 2.2: Conventional directed rewriting can go wrong in various ways if the wrong rewriting is applied at the wrong time. Orange highlighting indicates what changed from the initial term in each subfigure.

function symbols	f, g	
e-class ids	a, b	opaque identifiers
terms	$t ::= f \mid f(t_1, \dots, t_m)$	$m \geq 1$
e-nodes	$n ::= f \mid f(a_1, \dots, a_m)$	$m \geq 1$
e-classes	$c ::= \{n_1, \dots, n_m\}$	$m \geq 1$

Figure 2.3: Syntax and metavariables for the components of an e-graph. Function symbols may stand alone as constant e-nodes and terms. An e-class id is an opaque identifier that can be compared for equality with $=$.

Figure 2.2 shows some concrete examples of how poor rewrite choice can cause undesirable results.

2.2 E-Graphs

An *e-graph* is a data structure that stores a set of terms and a congruence relation over those terms. Originally developed for and still used in the heart of theorem provers [Nel80, DNS05, DMB08], e-graphs have also been used to power a program optimization technique called *equality saturation* [JNR02, TSTL09, STL11, NWA⁺20, PKSL20, WHL⁺20, PSSWT15].

2.2.1 Definitions

Intuitively, an e-graph is a set of equivalence classes (*e-classes*). Each e-class is a set of *e-nodes* representing equivalent terms from a given language, and an e-node is a function symbol paired with a list of children e-classes. More precisely:

Definition 2.1 (Definition of an E-Graph) *Given the definitions and syntax in Figure 2.3, an e-graph is a tuple (U, M, H) where:*

- A union-find data structure [Tar75] U stores an equivalence relation (denoted with $\equiv_{\text{id}}$) over e -class ids.
- The e -class map M maps e -class ids to e -classes. All equivalent e -class ids map to the same e -class, i.e., $a \equiv_{\text{id}} b$ iff $M[a]$ is the same set as $M[b]$. An e -class id a is said to refer to the e -class $M[\text{find}(a)]$.
- The hashcons¹ H is a map from e -nodes to e -class ids.

Note that an e -class has an identity (its canonical e -class id), but an e -node does not.² We use e -class id a and the e -class $M[\text{find}(a)]$ synonymously when clear from the context.

Definition 2.2 (Canonicalization) An e -graph's union-find U provides a find operation that canonicalizes e -class ids such that $\text{find}(U, a) = \text{find}(U, b)$ iff $a \equiv_{\text{id}} b$. We omit the first argument of find where clear from context.

- An e -class id a is canonical iff $\text{find}(a) = a$.
- An e -node n is canonical iff $n = \text{canonicalize}(n)$, where $\text{canonicalize}(f(a_1, a_2, \dots)) = f(\text{find}(a_1), \text{find}(a_2), \dots)$.

Definition 2.3 (Representation of Terms) An e -graph, e -class, or e -node is said to represent a term t if t can be “found” within it. Representation is defined recursively:

- An e -graph represents a term if any of its e -classes do.
- An e -class c represents a term if any e -node $n \in c$ does.

¹We use the term *hashcons* to evoke the memoization technique, since both avoid creating new duplicates of existing objects.

²Our definition of an e -graph reflects egg's design and therefore differs with some other e -graph definitions and implementations. In particular, making e -classes but not e -nodes identifiable is unique to our definition.

- An e-node $f(a_1, a_2, \dots)$ represents a term $f(t_1, t_2, \dots)$ if they have the same function symbol f and e-class $M[a_i]$ represents term t_i .

When each e-class is a singleton (containing only one e-node), an e-graph is essentially a term graph with sharing. Figure 2.4a shows an e-graph that represents the expression $(a \times 2)/2$.

Definition 2.4 (Equivalence) An e-graph defines three equivalence relations.

- Over e-class ids: $a \equiv_{\text{id}} b$ iff $\text{find}(a) = \text{find}(b)$.
- Over e-nodes: $n_1 \equiv_{\text{node}} n_2$ iff e-nodes n_1, n_2 are in the same e-class, i.e., $\exists a. n_1, n_2 \in M[a]$.
- Over terms: $t_1 \equiv_{\text{term}} t_2$ iff terms t_1, t_2 are represented in the same e-class.

We use $\equiv$ without the subscript when the relation is clear from context.

Definition 2.5 (Congruence) For a given e-graph, let $\cong$ denote a congruence relation over e-nodes such that $f(a_1, a_2, \dots) \cong f(b_1, b_2, \dots)$ iff $a_i \equiv_{\text{id}} b_i$. Let $\cong^*$ denote the congruence closure of $\equiv_{\text{node}}$, i.e., the smallest superset of $\equiv_{\text{node}}$ that is also a superset of $\cong$. Note that there may be two e-nodes such that $n_1 \cong^* n_2$ but $n_1 \not\equiv n_2$ and $n_1 \not\equiv_{\text{node}} n_2$. The relation $\cong$ only represents a single step of congruence; more than one step may be required to compute the congruence closure.

2.2.2 E-Graph Invariants

The e-graph must maintain invariants in order to correctly and efficiently implement the operations given in Section 2.2.3. This section only defines the invariants, discussion of how they are maintained is deferred to Chapter 3. These are collectively referred to as the *e-graph invariants*.

Definition 2.6 (The Congruence Invariant) The equivalence relation over e-nodes must be closed over congruence, i.e., $(\equiv_{\text{node}}) = (\cong^*)$. The e-graph must ensure that congruent e-nodes are in the

same *e*-class. Since identical *e*-nodes are trivially congruent, this implies that an *e*-node must be uniquely contained in a single *e*-class.

Definition 2.7 (The Hashcons Invariant) *The hashcons H must map all canonical *e*-nodes to their *e*-class ids. In other words:*

$$\text{e-node } n \in M[a] \iff H[\text{canonicalize}(n)] = \text{find}(a)$$

*If the hashcons invariant holds, then a procedure lookup can quickly find which *e*-class (if any) has an *e*-node congruent to a given *e*-node n : $\text{lookup}(n) = H[\text{canonicalize}(n)]$.*

2.2.3 Interface and Rewriting

E-graphs bear many similarities to the classic union-find data structure that they employ internally, and they inherit much of the terminology. E-graphs provide two main low-level mutating operations:

- **add** takes an *e*-node n and:
 - if $\text{lookup}(n) = a$, return a ;
 - if $\text{lookup}(n) = \emptyset$, then set $M[a] = \{n\}$ and return the id a .
- **merge** (sometimes called **assert** or **union**) takes two *e*-class ids a and b , unions them in the union-find U , and combines the *e*-classes by setting both $M[a]$ and $M[b]$ to $M[a] \cup M[b]$.

Both of these operations must take additional steps to maintain the congruence invariant. Invariant maintenance is discussed in Chapter 3.

E-graphs also offers operations for querying the data structure.

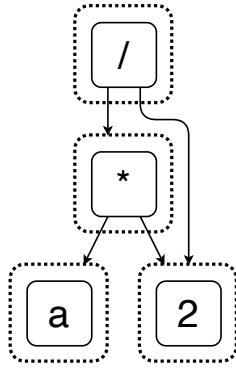
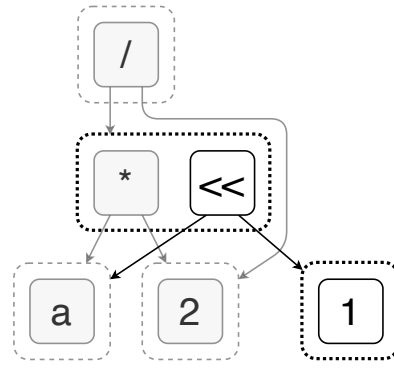
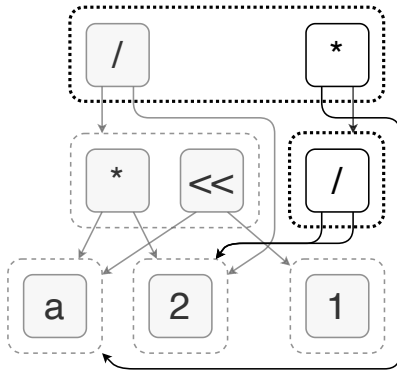
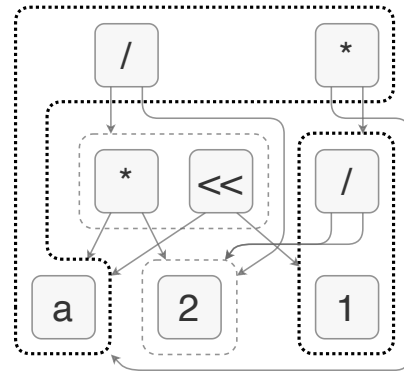
(a) Initial e-graph contains $(a \times 2)/2$.(b) Applied $x \times 2 \rightarrow x \ll 1$.(c) Applied rewrite $(x \times y)/z \rightarrow x \times (y/z)$.(d) Applied $x/x \rightarrow 1$ and $1 \times x \rightarrow x$.

Figure 2.4: An e-graph consists of e-classes (dashed boxes) containing equivalent e-nodes (solid boxes). Edges connect e-nodes to their child e-classes. Additions and modifications are emphasized in black. Applying rewrites to an e-graph adds new e-nodes and edges, but nothing is removed. Expressions added by rewrites are merged with the matched e-class. In Figure 2.4d, the rewrites do not add any new nodes, they only merge e-classes; so the e-graph gets *smaller* but represents *more* terms. Since the resulting e-graph has a cycle, it actually represents infinitely many expressions: a , $a \times 1$, $a \times 1 \times 1$, and so on.

- `find` canonicalizes e-class ids using the union-find U as described in definition 2.1.
- `ematch` performs the *e-matching* [DNS05, dMB07] procedure for finding patterns in the e-graph. `ematch` takes a pattern term p with variable placeholders and returns a list of tuples (σ, c) where σ is a substitution of variables to e-class ids such that $p[\sigma]$ is represented in e-class c .

These can be composed to perform rewriting over the e-graph. To apply a rewrite $\ell \rightarrow r$ to an e-graph, `ematch` finds tuples (σ, c) where e-class c represents $\ell[\sigma]$. Then, for each tuple, `merge( $c$ , add( $r[\sigma]$ ))` adds $r[\sigma]$ to the e-graph and unifies it with the matching e-class c .

Figure 2.4 shows an e-graph undergoing a series of rewrites. Note how the process is only additive; the initial term $(a \times 2)/2$ is still represented in the e-graph. Rewriting in an e-graph can also saturate, meaning the e-graph has learned every possible equivalence derivable from the given rewrites. If the user tried to apply $x \times y \rightarrow y \times x$ to an e-graph twice, the second time would add no additional e-nodes and perform no new merges; the e-graph can detect this and stop applying that rule.

2.3 Equality Saturation

Term rewriting [Der93] is a time-tested approach for equational reasoning in program optimization [TSTL09, JNR02], theorem proving [DNS05, DMB08], and program transformation [AEH⁺99]. In this setting, a tool repeatedly chooses one of a set of axiomatic rewrites, searches for matches of the left-hand pattern in the given expression, and replaces matching instances with the substituted right-hand side. It does, however, suffer from drawbacks such as the phase ordering problem (Section 2.1).

One solution to the phase ordering problem would simply apply all rewrites simultaneously,

```

1  def equality_saturation(expr, rewrites):
2      egraph = initial_egraph(expr)
3
4      while not egraph.is_saturated_or_timeout():
5
6          for rw in rewrites:
7              for (subst, eclass) in egraph.ematch(rw.lhs):
8                  eclass2 = egraph.add(rw.rhs.subst(subst))
9                  egraph.merge(eclass, eclass2)
10
11     return egraph.extract_best()

```

Figure 2.5: Pseudocode for equality saturation. Traditionally, equality saturation maintains the e-graph data structure invariants throughout the algorithm.

keeping track of every expression seen. This eliminates the problem of choosing the correct rule, but a naive implementation would require space exponential in the number of given rewrites. *Equality saturation* [TSTL09, STL11] is a technique to do this rewriting efficiently using an e-graph.

Figure 2.5 shows the equality saturation workflow. First, an initial e-graph is created from the input term. The core of the algorithm runs a set of rewrite rules until the e-graph is saturated (or a timeout is reached). Finally, a procedure called *extraction* selects the optimal represented term according to some cost function. For simple cost functions, a bottom-up, greedy traversal of the e-graph suffices to find the best term. Other extraction procedures have been explored for more complex cost functions [WHL⁺20, WZN⁺19].

Equality saturation eliminates the tedious and often error-prone task of choosing when to apply which rewrites, promising an appealingly simple workflow: state the relevant rewrites for the language, create an initial e-graph from a given expression, fire the rules until saturation, and finally extract the cheapest equivalent expression. Unfortunately, the technique remains ad hoc; prospective equality saturation users must implement their own e-graphs customized to their language, avoid performance pitfalls, and hack in the ability to do interpreted reasoning that is

not supported by purely syntactic rewrites. `egg` aims to address each aspect of these difficulties.

Equality Saturation and Theorem Proving An equality saturation engine and a theorem prover each have capabilities that would be impractical to replicate in the other. Automated theorem provers like satisfiability modulo theory (SMT) solvers are general tools that, in addition to supporting satisfiability queries, incorporate sophisticated, domain-specific solvers to allow interpreted reasoning within the supported theories. On the other hand, equality saturation is specialized for optimization, and its extraction procedure directly produces an optimal term with respect to a given cost function.

While SMT solvers are indeed the more general tool, equality saturation is not superseded by SMT; the specialized approach can be much faster when the full generality of SMT is not needed. To demonstrate this, we replicated a portion of the recent TASO paper [JPT⁺19], which optimizes deep learning models. As part of the work, they must verify a set of synthesized equalities with respect to a trusted set of universally quantified axioms. TASO uses Z3 [DMB08] to perform the verification even though most of Z3’s features (disjunctions, backtracking, theories, etc.) were not required. An equality saturation engine can also be used for verifying these equalities by adding the left and right sides of each equality to an e-graph, running the axioms as rewrites, and then checking if both sides end up in the same e-class. Z3 takes 24.65 seconds to perform the verification; `egg` performs the same task in 1.56 seconds (15× faster), or only 0.52 seconds (47× faster) when using `egg`’s batched evaluation (Section 5.3).

Equality Saturation and Superoptimization The Denali [JNR02] superoptimizer first demonstrated how to use e-graphs for optimized code generation as an alternative to hand-optimized machine code and prior exhaustive approaches [Mas87], both of which were less scalable. The inputs to Denali are programs in a C-like language from which it produces assembly programs.

Denali supported three types of rewrites—arithmetic, architectural, and program-specific. After applying these rewrites till saturation, it used architectural description of the hardware to generate constraints that were solved using a SAT solver to output a near-optimal program. While Denali’s approach was a significant improvement over prior work, it was intended to be used on straight line code only and therefore, did not apply to large real programs.

Equality saturation [TSTL09, STL11] developed a compiler optimization phase that works for complex language constructs like loops and conditionals. The first equality saturation paper used an intermediate representation called Program Expression Graphs (PEGs) to encode loops and conditionals. PEGs have specialized nodes that can represent infinite sequences, which allows them to represent loops. It uses a global profitability heuristic for extraction which is implemented using a pseudo-boolean solver. Recently, [PKSL20] used PEGs for code search. egg can support PEGs as a user-defined language, and thus their technique could be ported.

Chapter 3

Rebuilding: A New Take on E-graph Invariant Maintenance

Traditionally [Nel80, DNS05], e-graphs maintain their data structure invariants after each operation. We separate this invariant restoration into a procedure called *rebuilding*. This separation allows the client to choose when to enforce the e-graph invariants. Performing a rebuild immediately after every operation replicates the traditional approach to invariant maintenance. In contrast, rebuilding less frequently can amortize the cost of invariant maintenance, significantly improving performance.

In this section, we first describe how e-graphs have traditionally maintained invariants (Section 3.1). We then describe the rebuilding framework and how it captures a spectrum of invariant maintenance approaches, including the traditional one (Section 3.2). Using this flexibility, we then give a modified algorithm for equality saturation that enforces the e-graph invariants at only select points (Section 3.3). We finally demonstrate that this new approach offers an asymptotic speedup over traditional equality saturation (Section 3.4).

3.1 Upward Merging

Both mutating operations on the e-graph (add and merge, Section 2.2.3) can break the e-graph invariants if not done carefully. E-graphs have traditionally used *hashconsing* and *upward merging* to maintain the congruence invariant.

The add operation relies on the hashcons invariant (Definition 2.7) to quickly check whether the e-node n to be added—or one congruent to it—is already present. Without this check, add would create a new e-class with n in it even if some $n' \cong n$ was already in the e-graph, violating the congruence invariant.

The merge operation e-classes can violate both e-graph invariants. If $f(a, b)$ and $f(a, c)$ reside in two different e-classes x and y , merging b and c should also merge x and y to maintain the congruence invariant. This can propagate further, requiring additional merges.

E-graphs maintain a *parent list* for each e-class to maintain congruence. The parent list for e-class c holds all e-nodes that have c as a child. When merging two e-classes, e-graphs inspect these parent lists to find parents that are now congruent, recursively “upward merging” them if necessary.

The merge routine must also perform bookkeeping to preserve the hashcons invariant. In particular, merging two e-classes may change how parent e-nodes of those e-classes are canonicalized. The merge operation must therefore remove, re-canonicalize, and replace those e-nodes in the hashcons. In existing e-graph implementations [PSSWT15] used for equality saturation, maintaining the invariants while merging can take the vast majority of run time.

3.2 Rebuilding in Detail

Traditionally, invariant restoration is part of the merge operation itself. Rebuilding separates

```

1 def add(enode):
2     enode = self.canonicalize(enode)
3     if enode in self.hashcons:
4         return self.hashcons[enode]
5     else:
6         eclass_id = self.new_singleton_eclass(enode)
7         for child in enode.children:
8             child.parents.add(enode, eclass_id)
9         self.hashcons[enode] = eclass_id
10        return eclass_id
11
12 def merge(id1, id2):
13     if self.find(id1) == self.find(id2):
14         return self.find(id1)
15     new_id = self.union_find.union(id1, id2)
16     # traditional egraph merge can be
17     # emulated by calling rebuild right after
18     # adding the eclass to the worklist
19     self.worklist.add(new_id)
20     return new_id
21
22 def canonicalize(enode):
23     new_ch = [self.find(e) for e in enode.children]
24     return mk_enode(enode.op, new_ch)
25
26 def find(eclass_id):
27     return self.union_find.find(eclass_id)
28
29 def rebuild():
30     while self.worklist.len() > 0:
31         # empty the worklist into a local variable
32         todo = take(self.worklist)
33         # canonicalize and deduplicate the eclass refs
34         # to save calls to repair
35         todo = { self.find(eclass) for eclass in todo }
36         for eclass in todo:
37             self.repair(eclass)
38
39 def repair(eclass):
40     # update the hashcons so it always points
41     # canonical enodes to canonical eclasses
42     for (p_node, p_eclass) in eclass.parents:
43         self.hashcons.remove(p_node)
44         p_node = self.canonicalize(p_node)
45         self.hashcons[p_node] = self.find(p_eclass)
46
47     # deduplicate the parents, noting that equal
48     # parents get merged and put on the worklist
49     new_parents = {}
50     for (p_node, p_eclass) in eclass.parents:
51         p_node = self.canonicalize(p_node)
52         if p_node in new_parents:
53             self.merge(p_eclass, new_parents[p_node])
54         new_parents[p_node] = self.find(p_eclass)
55     eclass.parents = new_parents

```

Figure 3.1: Pseudocode for the add, merge, rebuild, and supporting methods. In each method, `self` refers to the e-graph being modified.

these concerns, reducing `merge`'s obligations and allowing for amortized invariant maintenance. In the rebuilding paradigm, `merge` maintains a *worklist* of e-class ids that need to be “upward merged”, i.e., e-classes whose parents are possibly congruent but not yet in the same e-class. The `rebuild` operation processes this worklist, restoring the invariants of deduplication and congruence. Rebuilding is similar to other approaches in how it restores congruence but it uniquely allows the client to choose when to restore invariants in the context of a larger algorithm like equality saturation.¹

Figure 3.1 shows pseudocode for the main e-graph operations and rebuilding. Note that `add` and `canonicalize` are given for completeness, but they are unchanged from the traditional e-graph implementation. The `merge` operation is similar, but it only adds the new e-class to the worklist instead of immediately starting upward merging. Adding a call to `rebuild` right after the addition to the worklist (Figure 3.1 line 19) would yield the traditional behavior of restoring the invariants immediately.

The `rebuild` method essentially calls `repair` on the e-classes from the worklist until the worklist is empty. Instead of directly manipulating the worklist, `egg`'s `rebuild` method first moves it into a local variable and deduplicates e-classes up to equivalence. Processing the worklist may merge e-classes, so breaking the worklist into chunks ensures that e-class ids made equivalent in the previous chunk are deduplicated in the subsequent chunk.

The actual work of `rebuild` occurs in the `repair` method. `repair` examines an e-class `c` and first canonicalizes e-nodes in the hashcons that have `c` as a child. Then it performs what is

¹Our rebuilding algorithm is similar to the congruence closure algorithm presented by [DST80]. The contribution of rebuilding is not *how* it restores the e-graph invariants but *when*; it gives the client the ability to specialize invariant restoration to a particular workload like equality saturation. Their algorithm also features a worklist of merges to be processed further, but it is offline, i.e., the algorithm processes a given set of equalities and outputs the set of equalities closed over congruence. Rebuilding is adapted to the online e-graph (and equality saturation) setting, where rewrites frequently examine the current set of equalities and assert new ones. Rebuilding additionally propagates e-class analysis facts (Section 4.1). Despite these differences, the core algorithms are similar enough that theoretical results on offline performance characteristics [DST80] apply to both. We do not provide theoretical analysis of rebuilding for the online setting; it is likely highly workload dependent.

essentially one “layer” of upward merging: if any of the parent e-nodes have become congruent, then their e-classes are merged and the result is added to the worklist.

Deduplicating the worklist, and thus reducing calls to `repair`, is at the heart of why deferring rebuilding improves performance. Intuitively, the upward merging process of rebuilding traces out a “path” of congruence through the e-graph. When rebuilding happens immediately after `merge` (and therefore frequently), these paths can substantially overlap. By deferring rebuilding, the chunk-and-deduplicate approach can coalesce the overlapping parts of these paths, saving what would have been redundant work. In our modified equality saturation algorithm (Section 3.3), deferred rebuilding is responsible for a significant, asymptotic speedup (Section 3.4).

3.2.1 Examples of Rebuilding

Deferred rebuilding speeds up congruence maintenance by amortizing the work of maintaining the hashcons invariant. Consider the following terms in an e-graph: $f_1(x), \dots, f_n(x), y_1, \dots, y_n$. Let the workload be $\text{merge}(x, y_1), \dots, \text{merge}(x, y_n)$. Each merge may change the canonical representation of the $f_i(x)$ s, so the traditional invariant maintenance strategy could require $O(n^2)$ hashcons updates. With deferred rebuilding the merges happen before the hashcons invariant is restored, requiring no more than $O(n)$ hashcons updates.

Deferred rebuilding can also reduce the number of calls to `repair`. Consider the following w terms in an e-graph, each nested under d function symbols:

$$f_1(f_2(\dots f_d(x_1))), \quad \dots, \quad f_1(f_2(\dots f_d(x_w)))$$

Note that w corresponds the width of this group of terms, and d to the depth. Let the workload be $w - 1$ merges that merge all the x s together: for $i \in [2, w]$, $\text{merge}(x_1, x_i)$.

In the traditional upward merging paradigm where `rebuild` is called after every `merge`, each

$\text{merge}(x_i, x_j)$ will require $O(d)$ calls to `repair` to maintain congruence, one for each layer of f_i . Over the whole workload, this requires $O(wd)$ calls to `repair`.

With deferred rebuilding, however, the $w - 1$ merges can all take place before congruence must be restored. Suppose the x s are all merged into an e-class c_x . When `rebuild` finally is called, the only element in the deduplicated worklist is c_x . Calling `repair` on c_x will merge the e-classes of the f_d e-nodes into an e-class c_{f_d} , adding the e-classes that contained those e-nodes back to the worklist. When the worklist is again deduplicated, c_{f_d} will be the only element, and the process repeats. Thus, the whole workload only incurs $O(d)$ calls to `repair`, eliminating the factor corresponding to the width of this group of terms. Figure 3.5 shows that the number calls to `repair` is correlated with time spent doing congruence maintenance.

3.2.2 Proof of Congruence

Intuitively, rebuilding is a delay of the upward merging process, allowing the user to choose when to restore the e-graph invariants. They are substantially similar in structure, with a critical difference in when the code is run. Below we offer a proof demonstrating that rebuilding restores the e-graph congruence invariant.

Theorem 3.1 *Rebuilding restores congruence and terminates.*

Proof 3.1 *Since rebuilding only merges congruent nodes, the congruence closure $\cong^*$ is fixed even though $\equiv_{\text{node}}$ changes. When $(\equiv_{\text{node}}) = (\cong^*)$, congruence is restored. Note that both $\equiv_{\text{node}}$ and $\cong^*$ are finite. We therefore show that rebuilding causes $\equiv_{\text{node}}$ to approach $\cong^*$. We define the set of incongruent e-node pairs as $I = (\cong^*) \setminus (\equiv_{\text{node}})$; in other words, $(n_1, n_2) \in I$ if $n_1 \cong^* n_2$ but $n_1 \not\equiv_{\text{node}} n_2$.*

Due to the additive nature of equality saturation, $\equiv_{\text{node}}$ only increases and therefore I is non-increasing. However, a call to `repair` inside the loop of `rebuild` does not necessarily shrink I . Some calls instead remove an element from the worklist but do not modify the e-graph at all.

Let the set W be the worklist of e -classes to be processed by `repair`; in Figure 3.1, W corresponds to `self.worklist` plus the unprocessed portion of the `todo` local variable. We show that each call to `repair` decreases the tuple $(|I|, |W|)$ lexicographically until $(|I|, |W|) = (0, 0)$, and thus rebuilding terminates with $(\equiv_{\text{node}}) = (\cong^*)$.

Given an e -class c from W , `repair` examines c 's parents for congruent e -nodes that are not yet in the same e -class:

- If at least one pair of c 's parents are congruent, rebuilding merges each pair (p_1, p_2) , which adds to W but makes I smaller by definition.
- If no such congruent pairs are found, do nothing. Then, $|W|$ is decreased by 1 since c came from the worklist and `repair` did not add anything back.

Since $(|I|, |W|)$ decreases lexicographically, $|W|$ eventually reaches 0, so `rebuild` terminates. Note that W contains precisely those e -classes that need to be “upward merged” to check for congruent parents. So, when W is empty, `rebuild` has effectively performed upward merging. By [Nel80, Chapter 7], $|I| = 0$. Therefore, when rebuilding terminates, congruence is restored.

3.3 Rebuilding and Equality Saturation

Rebuilding offers the choice of when to enforce the e -graph invariants, potentially saving work if deferred thanks to the deduplication of the worklist. The client is responsible for rebuilding at a time that maximizes performance without limiting the application.

`egg` provides a modified equality saturation algorithm to take advantage of rebuilding. Figure 3.2 shows pseudocode for both traditional equality saturation and `egg`'s variant, which exhibits two key differences:

```

1 def equality_saturation(expr, rewrites):
2     egraph = initial_egraph(expr)
3
4     while not egraph.is_saturated_or_timeout():
5
6         # reading and writing is mixed
7         for rw in rewrites:
8             for (subst, eclass) in egraph.ematch(rw.lhs):
9
10                # in traditional equality saturation,
11                # matches can be applied right away
12                # because invariants are always maintained
13                eclass2 = egraph.add(rw.rhs.subst(subst))
14                egraph.merge(eclass, eclass2)
15
16                # restore the invariants after each merge
17                egraph.rebuild()
18
19 return egraph.extract_best()
20

```

(a) Traditional equality saturation alternates between searching and applying rules, and the e-graph maintains its invariants throughout.

```

1 def equality_saturation(expr, rewrites):
2     egraph = initial_egraph(expr)
3
4     while not egraph.is_saturated_or_timeout():
5         matches = []
6
7         # read-only phase, invariants are preserved
8         for rw in rewrites:
9             for (subst, eclass) in egraph.ematch(rw.lhs):
10                matches.append((rw, subst, eclass))
11
12        # write-only phase, temporarily break invariants
13        for (rw, subst, eclass) in matches:
14            eclass2 = egraph.add(rw.rhs.subst(subst))
15            egraph.merge(eclass, eclass2)
16
17        # restore the invariants once per iteration
18        egraph.rebuild()
19
20 return egraph.extract_best()
21

```

(b) egg splits equality saturation iterations into read and write phases. The e-graph invariants are not constantly maintained, but restored only at the end of each iteration by the rebuild method (Chapter 3).

Figure 3.2: Pseudocode for traditional and egg’s version of the equality saturation algorithm.

1. Each iteration is split into a read phase, which searches for all the rewrite matches, and a write phase that applies those matches.²
2. Rebuilding occurs only once per iteration, at the end.

egg’s separation of the read and write phases means that rewrites are truly unordered. In traditional equality saturation, later rewrites in the given rewrite list are favored in the sense that they can “see” the results of earlier rewrites in the same iteration. Therefore, the results depend on the order of the rewrite list if saturation is not reached (which is common on large rewrite lists or input expressions). egg’s equality saturation algorithm is invariant to the order of the rewrite list.

Separating the read and write phases also allows egg to safely defer rebuilding. If rebuilding were deferred in the traditional equality saturation algorithm, rules later in the rewrite list would be searched against an e-graph with broken invariants. Since congruence may not hold, there may be missing equivalences, resulting in missing matches. These matches will be seen after the `rebuild` during the next iteration (if another iteration occurs), but the false reporting could impact metrics collection, rule scheduling,³ or saturation detection.

3.4 Evaluating Rebuilding

To demonstrate that deferred rebuilding provides faster congruence closure than traditional upward merging, we modified egg to call `rebuild` immediately after every merge. This provides a one-to-one comparison of deferred rebuilding against the traditional approach, isolated from the many other factors that make egg efficient: overall design and algorithmic differences, programming

²Although the original equality saturation paper [TSTL09] does not have separate reading and writing phases, some e-graph implementations (like the one inside Z3 [DMB08]) do separate these phases as an implementation detail. Ours is the first algorithm to take advantage of this by deferring invariant maintenance.

³An optimization introduced in Figure 5.2 that relies on an accurate count of how many times a rewrite was matched.

language performance, and other orthogonal performance improvements.

We ran `egg`'s test suite using both rebuild strategies, measuring the time spent on congruence maintenance. Each test consists of one run of `egg`'s equality saturation algorithm to optimize a given expression. Of the 32 total tests, 8 hit the iteration limit of 100 and the remainder saturated. Note that both rebuilding strategies use `egg`'s phase-split equality saturation algorithm, and the resulting e-graphs are identical in all cases. These experiments were performed on a 2020 Macbook Pro with a 2 GHz quad-core Intel Core i5 processor and 16GB of memory.

Figure 3.3 shows our how rebuilding speeds up congruence maintenance. Overall, our experiments show an aggregate $87.85\times$ speedup on congruence closure and $20.96\times$ speedup over the entire equality saturation algorithm. Figure 3.4 shows this speedup is asymptotic; the multiplicative speedup increases as problem gets larger.

`egg`'s test suite consists of two main applications: `math`, a small computer algebra system capable of symbolic differentiation and integration; and `lambda`, a partial evaluator for the untyped lambda calculus using explicit substitution to handle variable binding (shown in Chapter 5). Both are typical `egg` applications primarily driven by syntactic rewrites, with a few key uses of `egg`'s more complex features like e-class analyses and dynamic/conditional rewrites.

`egg` can be configured to capture various metrics about equality saturation as it runs, including the time spent in the read phase (searching for matches), the write phase (applying matches), and rebuilding. In Figure 3.3, congruence time is measured as the time spent applying matches plus rebuilding. Other parts of the equality saturation algorithm (creating the initial e-graph, extracting the final term) take negligible take compared to the equality saturation iterations.

Deferred rebuilding amortizes the examination of e-classes for congruence maintenance; deduplicating the worklist reduces the number of calls to the `repair`. Figure 3.5 shows that time spent in congruence is correlated with the number of calls to the `repair` methods.

The case study in Section 6.2 provides a further evaluation of rebuilding. Rebuilding (and

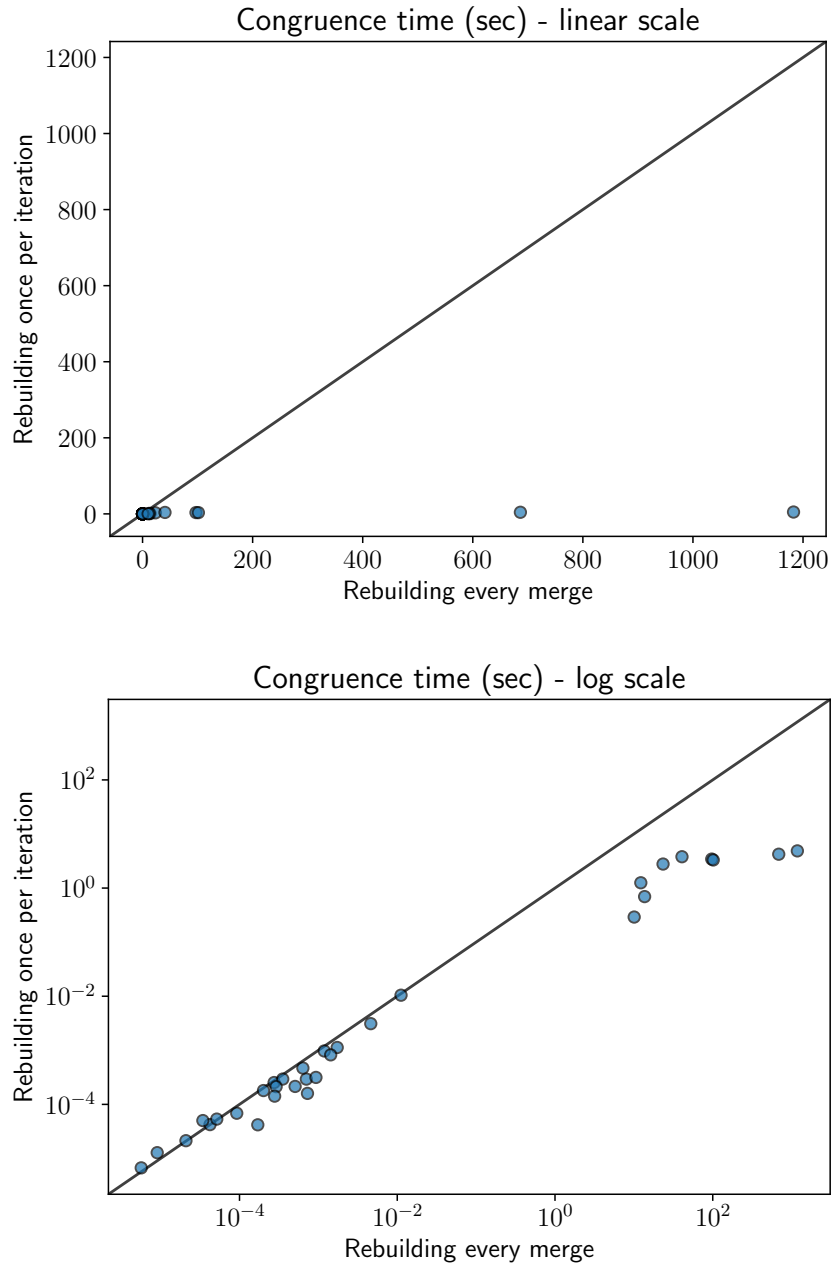


Figure 3.3: Rebuilding once per iteration—as opposed to after every merge—significantly speeds up congruence maintenance. Both plots show the same data: one point for each of the 32 tests. The diagonal line is $y = x$; points below the line mean deferring rebuilding is faster. In aggregate over all tests (using geometric mean), congruence is 87.85 $\times$ faster, and equality saturation is 20.96 $\times$ faster. The linear scale plot shows that deferred rebuilding is significantly faster. The log scale plot suggests the speedup is greater than some constant multiple; Figure 3.4 demonstrates this in greater detail.

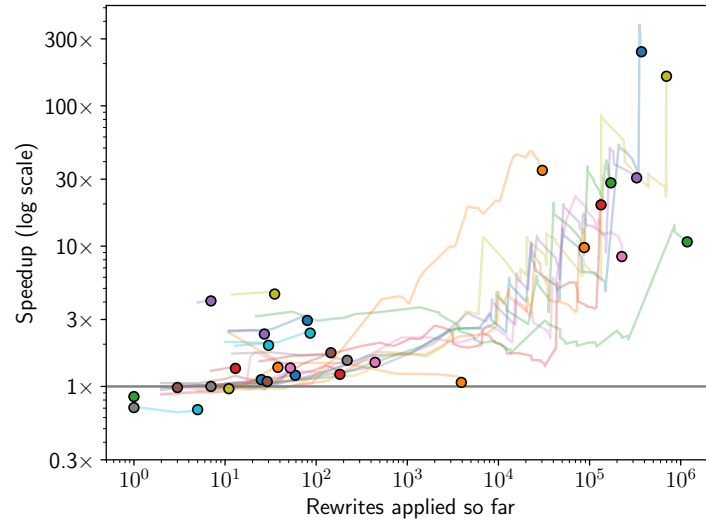


Figure 3.4: As more rewrites are applied, deferring rebuilding gives greater speedup. Each line represents a single test: each equality saturation iteration plots the cumulative rewrites applied so far against the multiplicative speedup of deferring rebuilding; the dot represents the end of that test. Both the test suite as a whole (the dots) and individual tests (the lines) demonstrate an asymptotic speedup that increases with the problem size.

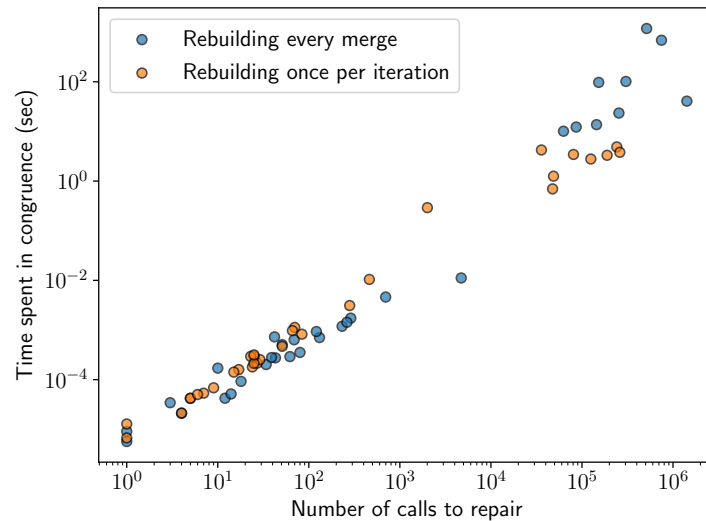


Figure 3.5: The time spent in congruence maintenance correlates with the number of calls to the repair method. Spearman correlation yields $r = 0.98$ with a p-value of $3.6e-47$, indicating that the two quantities are indeed positively correlated.

other egg features) have also been implemented in a Racket-based e-graph, demonstrating that rebuilding is a conceptual advance that need not be tied to the egg implementation.

Chapter 4

Extending E-graphs with E-class Analyses

As discussed so far, e-graphs and equality saturation provide an efficient way to implement a term rewriting system. Rebuilding enhances that efficiency, but the approach remains designed for purely syntactic rewrites. However, program analysis and optimization typically require more than just syntactic information. Instead, transformations are *computed* based on the input terms and also semantic facts about that input term, e.g., constant value, free variables, nullability, numerical sign, size in memory, and so on. The “purely syntactic” restriction has forced existing equality saturation applications [TSTL09, STL11, PSSWT15] to resort to ad hoc passes over the e-graph to implement analyses like constant folding. These ad hoc passes require manually manipulating the e-graph, the complexity of which could prevent the implementation of more sophisticated analyses.

We present a new technique called *e-class analysis*, which allows the concise expression of a program analysis over the e-graph. An e-class analysis resembles abstract interpretation lifted to the e-graph level, attaching *analysis data* from a semilattice to each e-class. The e-graph maintains and propagates this data as e-classes get merged and new e-nodes are added. Analysis data can be used directly to modify the e-graph, to inform how or if rewrites apply their right-hand sides,

or to determine the cost of terms during the extraction process.

E-class analyses provide a general mechanism to replace what previously required ad hoc extensions that manually manipulate the e-graph. E-class analyses also fit within the equality saturation workflow, so they can naturally cooperate with the equational reasoning provided by rewrites. Moreover, an analysis lifted to the e-graph level automatically benefits from a sort of “partial-order reduction” for free: large numbers of similar programs may be analyzed for little additional cost thanks to the e-graph’s compact representation.

This section provides a conceptual explanation of e-class analyses as well as dynamic and conditional rewrites that can use the analysis data. The following sections will provide concrete examples: Chapter 5 discusses the egg implementation and a complete example of a partial evaluator for the lambda calculus; Chapter 6 discusses how three published projects have used egg and its unique features (like e-class analyses).

4.1 E-Class Analyses

An e-class analysis defines a domain D and associates a value $d_c \in D$ to each e-class c . The e-class c contains the associated data d_c , i.e., given an e-class c , one can get d_c easily, but not vice-versa.

The interface of an e-class analysis is as follows, where G refers to the e-graph, and n and c refer to e-nodes and e-classes within G :

- $\text{make}(n) \rightarrow d_c$ When a new e-node n is added to G into a new, singleton e-class c , construct a new value $d_c \in D$ to be associated with n 's new e-class, typically by accessing the associated data of n 's children.
- $\text{join}(d_{c_1}, d_{c_2}) \rightarrow d_c$ When e-classes c_1, c_2 are being merged into c , join d_{c_1}, d_{c_2} into a new value d_c to be associated with the new e-class c .
- $\text{modify}(c) \rightarrow c'$ Optionally modify the e-class c based on d_c , typically by adding an e-node to c . Modify should be idempotent if no other changes occur to the e-class, i.e., $\text{modify}(\text{modify}(c)) = \text{modify}(c)$

The domain D together with the join operation should form a join-semilattice. The semilattice perspective is useful for defining the *analysis invariant* (where $\vee$ is the join operation):

$$\forall c \in G. \quad d_c = \bigvee_{n \in c} \text{make}(n) \quad \text{and} \quad \text{modify}(c) = c$$

The first part of the analysis invariant states that the data associated with each e-class must be the join of the make for every e-node in that e-class. Since D is a join-semilattice, this means that $\forall c, \forall n \in c, d_c \geq \text{make}(n)$. The motivation for the second part is more subtle. Since the analysis can modify an e-class through the modify method, the analysis invariant asserts that these modifications are driven to a fixed point. When the analysis invariant holds, a client looking at the analysis data can be assured that the analysis is “stable” in the sense that recomputing make, join, and modify will not modify the e-graph or any analysis data.

4.1.1 Maintaining the Analysis Invariant

We extend the rebuilding procedure from Chapter 3 to restore the analysis invariant as well as the congruence invariant. Figure 4.1 shows the necessary modifications to the rebuilding code

```

1 def add(enode):
2     enode = self.canonicalize(enode)
3     if enode in self.hashcons:
4         return self.hashcons[enode]
5     else:
6         eclass = self.new_singleton_eclass(enode)
7         for child_eclass in enode.children:
8             child_eclass.parents.add(enode, eclass)
9         self.hashcons[enode] = eclass
10        eclass.data = analysis.make(enode)
11        analysis.modify(eclass)
12        return eclass
13
14 def merge(eclass1, eclass2)
15     union = self.union_find.union(eclass1, eclass2)
16     if not union.was_already_unioned:
17         d1, d2 = eclass1.data, eclass2.data
18         union.eclass.data = analysis.join(d1, d2)
19         self.worklist.add(union.eclass)
20     return union.eclass
21
22 def repair(eclass):
23     for (p_node, p_eclass) in eclass.parents:
24         self.hashcons.remove(p_node)
25         p_node = self.canonicalize(p_node)
26         self.hashcons[p_node] = self.find(p_eclass)
27
28     new_parents = {}
29     for (p_node, p_eclass) in eclass.parents:
30         p_node = self.canonicalize(p_node)
31         if p_node in new_parents:
32             self.union(p_eclass, new_parents[p_node])
33         new_parents[p_node] = self.find(p_eclass)
34     eclass.parents = new_parents
35
36     # any mutations modify makes to eclass
37     # will add to the worklist
38     analysis.modify(eclass)
39     for (p_node, p_eclass) in eclass.parents:
40         new_data = analysis.join(
41             p_eclass.data,
42             analysis.make(p_node))
43         if new_data != p_eclass.data:
44             p_eclass.data = new_data
45             self.worklist.add(p_eclass)

```

Figure 4.1: The pseudocode for maintaining the e-class analysis invariant is largely similar to how rebuilding maintains congruence closure (Chapter 3). Only lines 10–11, 17–18, and 37–44 are added. Grayed out or missing code is unchanged from Figure 3.1.

from Figure 3.1.

Adding e-nodes and merging e-classes risk breaking the analysis invariant in different ways. Adding e-nodes is the simpler case; lines 10–11 restore the invariant for the newly created, singleton e-class that holds the new e-node. When merging e-nodes, the first concern is maintaining the semilattice portion of the analysis invariant. Since join forms a semilattice over the domain D of the analysis data, the order in which the joins occur does not matter. Therefore, line 18 suffices to update the analysis data of the merged e-class.

Since `make( $n$ )` creates analysis data by looking at the data of n 's, children, merging e-classes can violate the analysis invariant in the same way it can violate the congruence invariant. The solution is to use the same worklist mechanism introduced in Chapter 3. Lines 37–44 of the `repair` method (which `rebuild` on each element of the worklist) re-make and merge the analysis data of the parent of any recently merged e-classes. The new `repair` method also calls `modify` once, which suffices due to its idempotence. In the pseudocode, `modify` is reframed as a mutating method for clarity.

`egg`'s implementation of e-class analyses assumes that the analysis domain D is indeed a semilattice and that `modify` is idempotent. Without these properties, `egg` may fail to restore the analysis invariant on `rebuild`, or it may not terminate.

4.1.2 Example: Constant Folding

The data produced by e-class analyses can be usefully consumed by other components of an equality saturation system (see Section 4.2), but e-class analyses can be useful on their own thanks to the `modify` hook. Typical `modify` hooks will either do nothing, check some invariant about the e-classes being merged, or add an e-node to that e-class (using the regular `add` and `merge` methods of the e-graph).

As mentioned above, other equality saturation implementations have implemented constant folding as custom, ad hoc passes over the e-graph. We can formulate constant folding as an e-class analysis that highlights the parallels with abstract interpretation. Let the domain $D = \text{Option}\langle\text{Constant}\rangle$, and let the join operation be the “or” operation of the Option type: Note

```

match (a, b) {
  (None,    None    ) => None,
  (Some(x), None    ) => Some(x),
  (None,    Some(y)) => Some(y),
  (Some(x), Some(y)) => { assert!(x == y); Some(x) }
}

```

how join can also aid in debugging by checking properties about values that are unified in the e-graph; in this case we assert that all terms represented in an e-class should have the same constant value. The make operation serves as the abstraction function, returning the constant value of an e-node if it can be computed from the constant values associated with its children e-classes. The modify operation serves as a concretization function in this setting. If d_c is a constant value, then $\text{modify}(c)$ would add $\gamma(d_c) = n$ to c , where γ concretizes the constant value into a childless e-node.

Constant folding is an admittedly simple analysis, but one that did not formerly fit within the equality saturation framework. E-class analyses support more complicated analyses in a general way, as discussed in later sections on the egg implementation and case studies (Sections 5 and 6).

4.2 Conditional and Dynamic Rewrites

In equality saturation applications, most of the rewrites are purely syntactic. In some cases, additional data may be needed to determine if or how to perform the rewrite. For example, the rewrite $x/x \rightarrow 1$ is only valid if $x \neq 0$. A more complex rewrite may need to compute the

right-hand side dynamically based on an analysis fact from the left-hand side.

The right-hand side of a rewrite can be generalized to a function apply that takes a substitution and an e-class generated from e-matching the left-hand side, and produces a term to be added to the e-graph and unified with the matched e-class. For a purely syntactic rewrite, the apply function need not inspect the matched e-class in any way; it would simply apply the substitution to the right-hand pattern to produce a new term.

E-class analyses greatly increase the utility of this generalized form of rewriting. The apply function can look at the analysis data for the matched e-class or any of the e-classes in the substitution to determine if or how to construct the right-hand side term. These kinds of rewrites can be broken down further into two categories:

- *Conditional* rewrites like $x/x \rightarrow 1$ that are purely syntactic but whose validity depends on checking some analysis data;
- *Dynamic* rewrites that compute the right-hand side based on analysis data.

Conditional rewrites are a subset of the more general dynamic rewrites. Our egg implementation supports both. The example in Chapter 5 and case studies in Chapter 6 heavily use generalized rewrites, as it is typically the most convenient way to incorporate domain knowledge into the equality saturation framework.

4.3 Extraction

Equality saturation typically ends with an extraction phase that selects an optimal represented term from an e-class according to some cost function. In many domains [PSSWT15, NWA⁺20], AST size (sometimes weighted differently for different operators) suffices as a simple, local cost function. We say a cost function k is local when the cost of a term $f(a_1, \dots)$ can be computed

from the function symbol f and the costs of the children. With such cost functions, extracting an optimal term can be efficiently done with a fixed-point traversal over the e-graph that selects the minimum cost e-node from each e-class [PSSWT15].

Extraction can be formulated as an e-class analysis when the cost function is local. The analysis data is a tuple $(n, k(n))$ where n is the cheapest e-node in that e-class and $k(n)$ its cost. The $\text{make}(n)$ operation calculates the cost $k(n)$ based on the analysis data (which contain the minimum costs) of n 's children. The merge operation simply takes the tuple with lower cost. The semilattice portion of the analysis invariant then guarantees that the analysis data will contain the lowest-cost e-node in each class. Extract can then proceed recursively; if the analysis data for e-class c gives $f(c_1, c_2, \dots)$ as the optimal e-node, the optimal term represented in c is $\text{extract}(c) = f(\text{extract}(c_1), \text{extract}(c_2), \dots)$. This not only further demonstrates the generality of e-class analyses, but also provides the ability to do extraction “on the fly”; conditional and dynamic rewrites can determine their behavior based on the cheapest term in an e-class.

Extraction (whether done as a separate pass or an e-class analysis) can also benefit from the analysis data. Typically, a local cost function can only look at the function symbol of the e-node n and the costs of n 's children. When an e-class analysis is attached to the e-graph, however, a cost function may observe the data associated with n 's e-class, as well as the data associated with n 's children. This allows a cost function to depend on computed facts rather than just purely syntactic information. In other words, the cost of an operator may differ based on its inputs. Section 6.3 provides a motivating case study wherein an e-class analysis computes the size and shape of tensors, and this size information informs the cost function.

Chapter 5

egg: Easy, Extensible, and Efficient E-graphs

We implemented the techniques of rebuilding and e-class analysis in `egg`, an easy-to-use, extensible, and efficient e-graph library. To the best of our knowledge, `egg` is the first general-purpose, reusable e-graph implementation. This has allowed focused effort on ease of use and optimization, knowing that any benefits will be seen across use cases as opposed to a single, ad hoc instance.

This section details `egg`'s implementation and some of the various optimizations and tools it provides to the user. We use an extended example of a partial evaluator for the lambda calculus¹, for which we provide the complete source code (which few changes for readability) in Figure 5.1 and Figure 5.2. While contrived, this example is compact and familiar, and it highlights (1) how `egg` is used and (2) some of its novel features like e-class analyses and dynamic rewrites. It demonstrates how `egg` can tackle binding, a perennially tough problem for e-graphs, with a simple explicit substitution approach powered by `egg`'s extensibility. Chapter 6 goes further, providing

¹E-graphs do not have any “built-in” support for binding; for example, equality modulo alpha renaming is not free. The explicit substitution provided in this section is illustrative but rather high in performance cost. Better support for languages with binding is important future work.

real-world case studies of published projects that have depended on egg.

egg is implemented in ~5000 lines of Rust,² including code, tests, and documentation. egg is open-source, well-documented, and distributed via Rust’s package management system.³ All of egg’s components are generic over the user-provided language, analysis, and cost functions.

5.1 Ease of Use

egg’s ease of use comes primarily from its design as a library. By defining only a language and some rewrite rules, a user can quickly start developing a synthesis or optimization tool. Using egg as a Rust library, the user defines the language using the `define_language!` macro shown in Figure 5.1, lines 1-22. Childless variants in the language may contain data of user-defined types, and e-class analyses or dynamic rewrites may inspect this data.

The user provides rewrites as shown in Figure 5.1, lines 51-100. Each rewrite has a name, a left-hand side, and a right-hand side. For purely syntactic rewrites, the right-hand is simply a pattern. More complex rewrites can incorporate conditions or even dynamic right-hand sides, both explained in the Section 5.2 and Figure 5.2.

Equality saturation workflows, regardless of the application domain, typically have a similar structure: add expressions to an empty e-graph, run rewrites until saturation or timeout, and extract the best equivalent expressions according to some cost function. This “outer loop” of equality saturation involves a significant amount of error-prone boilerplate:

- Checking for saturation, timeouts, and e-graph size limits.

²Rust [Rus] is a high-level systems programming language. egg has been integrated into applications written in other programming languages using both C FFI and serialization approaches.

³Source: <https://github.com/mwillsey/egg>. Documentation: <https://docs.rs/egg>. Package: <https://crates.io/crates/egg>. This paper uses version 0.6 of egg.

```

1  define_language! {
2      enum Lambda {
3          // enum variants have data or children (e-class Ids)
4          // [Id; N] is an array of N 'Id's
5
6          // base type operators
7          "+" = Add([Id; 2]), "=" = Eq([Id; 2]),
8          "if" = If([Id; 3]),
9
10         // functions and binding
11         "app" = App([Id; 2]), "lam" = Lambda([Id; 2]),
12         "let" = Let([Id; 3]), "fix" = Fix([Id; 2]),
13
14         // (var x) is a use of 'x' as an expression
15         "var" = Use(Id),
16         // (subst a x b) substitutes a for (var x) in b
17         "subst" = Subst([Id; 3]),
18
19         // base types have no children, only data
20         Bool(bool), Num(i32), Symbol(String),
21     }
22 }
23
24 // example terms and what they simplify to
25 // pulled directly from the egg test suite
26
27 test_fn! { lambda_under, rules(),
28     "(lam x (+ 4 (app (lam y (var y)) 4)))"
29     => "(lam x 8)",
30 }
31
32 test_fn! { lambda_compose_many, rules(),
33     "(let compose (lam f (lam g (lam x
34         (app (var f)
35             (app (var g) (var x))))))
36     (let add1 (lam y (+ (var y) 1))
37     (app (app (var compose) (var add1))
38         (app (app (var compose) (var add1))
39             (app (app (var compose) (var add1))
40                 (app (app (var compose) (var add1))
41                     (var add1))))))"
42     => "(lam ?x (+ (var ?x) 5))"
43 }
44
45 test_fn! { lambda_if_elim, rules(),
46     "(if (= (var a) (var b))
47         (+ (var a) (var a))
48         (+ (var a) (var b)))"
49     => "(+ (var a) (var b))"
50 }
51
52 // Returns a list of rewrite rules
53 fn rules() -> Vec<Rewrite<Lambda, LambdaAnalysis>> { vec![
54     // open term rules
55     rw!("if-true"; "(if true ?then ?else)" => "?then"),
56     rw!("if-false"; "(if false ?then ?else)" => "?else"),
57     rw!("if-elim"; "(if (= (var ?x) ?e) ?then ?else)" => "?else"
58         if ConditionEqual::parse("(let ?x ?e ?then)",
59             "(let ?x ?e ?else)")),
60     rw!("add-comm"; "(+ ?a ?b)" => "(+ ?b ?a)"),
61     rw!("add-assoc"; "(+ (+ ?a ?b) ?c)" => "(+ ?a (+ ?b ?c))"),
62     rw!("eq-comm"; "(= ?a ?b)" => "(= ?b ?a)"),
63
64     // substitution introduction
65     rw!("fix"; "(fix ?v ?e)" =>
66         "(let ?v (fix ?v ?e) ?e)"),
67     rw!("beta"; "(app (lam ?v ?body) ?e)" =>
68         "(let ?v ?e ?body)"),
69
70     // substitution propagation
71     rw!("let-app"; "(let ?v ?e (app ?a ?b))" =>
72         "(app (let ?v ?e ?a) (let ?v ?e ?b))"),
73     rw!("let-add"; "(let ?v ?e (+ ?a ?b))" =>
74         "(+ (let ?v ?e ?a) (let ?v ?e ?b))"),
75     rw!("let-eq"; "(let ?v ?e (= ?a ?b))" =>
76         "(= (let ?v ?e ?a) (let ?v ?e ?b))"),
77     rw!("let-if"; "(let ?v ?e (if ?cond ?then ?else))" =>
78         "(if (let ?v ?e ?cond)
79             (let ?v ?e ?then)
80             (let ?v ?e ?else))"),
81
82     // substitution elimination
83     rw!("let-const"; "(let ?v ?e ?c)" => "?c"
84         if is_const(var("?c"))),
85     rw!("let-var-same"; "(let ?v1 ?e (var ?v1))" => "?e"),
86     rw!("let-var-diff"; "(let ?v1 ?e (var ?v2))" => "(var ?v2)"
87         if is_not_same_var(var("?v1"), var("?v2"))),
88     rw!("let-lam-same"; "(let ?v1 ?e (lam ?v1 ?body))" =>
89         "(lam ?v1 ?body)"),
90     rw!("let-lam-diff"; "(let ?v1 ?e (lam ?v2 ?body))" =>
91         (CaptureAvoid {
92             fresh: var("?fresh"), v2: var("?v2"), e: var("?e"),
93             if_not_free: "(lam ?v2 (let ?v1 ?e ?body))"
94                 .parse().unwrap(),
95             if_free: "(lam ?fresh (let ?v1 ?e
96                 (let ?v2 (let ?v1 ?e ?fresh) ?body)))"
97                 .parse().unwrap(),
98         })
99         if is_not_same_var(var("?v1"), var("?v2"))),
100 }

```

Figure 5.1: egg is generic over user-defined languages; here we define a language and rewrite rules for a lambda calculus partial evaluator. The provided `define_language!` macro (lines 1-22) allows the simple definition of a language as a Rust enum, automatically deriving parsing and pretty printing. A value of type `Lambda` is an e-node that holds either data that the user can inspect or some number of e-class children (e-class Ids).

Rewrite rules can also be defined succinctly (lines 51-100). Patterns are parsed as s-expressions: strings from the `define_language!` invocation (ex: `fix`, `=`, `+`) and data from the variants (ex: `false`, `1`) parse as operators or terms; names prefixed by “?” parse as pattern variables.

Some of the rewrites made are conditional using the “left => right if cond” syntax. The `if-elim` rewrite on line 57 uses egg’s provided `ConditionEqual` as a condition, only applying the right-hand side if the e-graph can prove the two argument patterns equivalent. The final rewrite, `let-lam-diff`, is dynamic to support capture avoidance; the right-hand side is a Rust value that implements the `Applier` trait instead of a pattern. Figure 5.2 contains the supporting code for these rewrites.

We also show some of the tests (lines 27-50) from egg’s lambda test suite. The tests proceed by inserting the term on the left-hand side, running egg’s equality saturation, and then checking to make sure the right-hand pattern can be found in the same e-class as the initial term.

- Orchestrating the read-phase, write-phase, rebuild system (Figure 3.1) that makes egg fast.
- Recording performance data at each iteration.
- Potentially coordinating rule execution so that expansive rules like associativity do not dominate the e-graph.
- Finally, extracting the best expression(s) according to a user-defined cost function.

egg provides these functionalities through its `Runner` and `Extractor` interfaces. Runners automatically detect saturation, and can be configured to stop after a time, e-graph size, or iterations limit. The equality saturation loop provided by egg calls `rebuild`, so users need not even know about egg’s deferred invariant maintenance. Runners record various metrics about each iteration automatically, and the user can hook into this to report relevant data. Extractors select the optimal term from an e-graph given a user-defined, local cost function.⁴ The two can be combined as well; users commonly record the “best so far” expression by extracting in each iteration.

Figure 5.1 also shows egg’s `test_fn!` macro for easily creating tests (lines 27-50). These tests create an e-graph with the given expression, run equality saturation using a `Runner`, and check to make sure the right-hand pattern can be found in the same e-class as the initial expression.

5.2 Extensibility

For simple domains, defining a language and purely syntactic rewrites will suffice. However, our partial evaluator requires interpreted reasoning, so we use some of egg’s more advanced features

⁴As mentioned in Section 4.3, extraction can be implemented as part of an e-class analysis. The separate `Extractor` feature is still useful for ergonomic and performance reasons.

like e-class analyses and dynamic rewrites. Importantly, `egg` supports these extensibility features as a library: the user need not modify the e-graph or `egg`'s internals.

Figure 5.2 shows the remainder of the code for our lambda calculus partial evaluator. It uses an e-class analysis (`LambdaAnalysis`) to track free variables and constants associated with each e-class. The implementation of the e-class analysis is in Lines 11-50. The e-class analysis invariant guarantees that the analysis data contains an over-approximation of free variables from terms represented in that e-class. The analysis also does constant folding (see the `make` and `modify` methods). The `let-lam-diff` rewrite (Line 90, Figure 5.1) uses the `CaptureAvoid` (Lines 81-100, Figure 5.2) dynamic right-hand side to do capture-avoiding substitution only when necessary based on the free variable information. The conditional rewrites from Figure 5.1 depend on the conditions `is_not_same_var` and `is_var` (Lines 68-74, Figure 5.2) to ensure correct substitution.

`egg` is extensible in other ways as well. As mentioned above, `Extractors` are parameterized by a user-provided cost function. `Runners` are also extensible with user-provided rule schedulers that can control the behavior of potentially troublesome rewrites. In typical equality saturation, each rewrite is searched for and applied each iteration. This can cause certain rewrites, commonly associativity or distributivity, to dominate others and make the search space less productive. Applied in moderation, these rewrites can trigger other rewrites and find greatly improved expressions, but they can also slow the search by exploding the e-graph exponentially in size. By default, `egg` uses the built-in backoff scheduler that identifies rewrites that are matching in exponentially-growing locations and temporarily bans them. We have observed that this greatly reduced run time (producing the same results) in many settings. `egg` can also use a conventional every-rule-every-time scheduler, or the user can supply their own.

```

1 type EGraph = egg::EGraph<Lambda, LambdaAnalysis>;
2 struct LambdaAnalysis;
3 struct FC {
4   free: HashSet<Id>, // our analysis data stores free vars
5   constant: Option<Lambda>, // and the constant value, if any
6 }
7
8 // helper function to make pattern meta-variables
9 fn var(s: &str) -> Var { s.parse().unwrap() }
10
11 impl Analysis<Lambda> for LambdaAnalysis {
12   type Data = FC; // attach an FC to each eclass
13   // merge implements semilattice join by joining into 'to'
14   // returning true if the 'to' data was modified
15   fn merge(&self, to: &mut FC, from: FC) -> bool {
16     let before_len = to.free.len();
17     // union the free variables
18     to.free.extend(from.free.iter().copied());
19     if to.constant.is_none() && from.constant.is_some() {
20       to.constant = from.constant;
21     } else {
22       before_len != to.free.len()
23     }
24   }
25 }
26
27 fn make(egraph: &EGraph, enode: &Lambda) -> FC {
28   let f = |i: &Id| egraph[*i].data.free.iter().copied();
29   let mut free = HashSet::default();
30   match enode {
31     Use(v) => { free.insert(*v); }
32     Let([v, a, b]) => {
33       free.extend(f(b)); free.remove(v); free.extend(f(a));
34     }
35     Lambda([v, b]) | Fix([v, b]) => {
36       free.extend(f(b)); free.remove(v);
37     }
38     _ => enode.for_each_child(
39       |c| free.extend(&egraph[c].data.free),
40     )
41   }
42   FC { free: free, constant: eval(egraph, enode) }
43 }
44
45 fn modify(egraph: &mut EGraph, id: Id) {
46   if let Some(c) = egraph[id].data.constant.clone() {
47     let const_id = egraph.add(c);
48     egraph.union(id, const_id);
49   }
50 }
51
52 // evaluate an enode if the children have constants
53 // Rust's '?' extracts an Option, early returning if None
54 fn eval(eg: &EGraph, enode: &Lambda) -> Option<Lambda> {
55   let c = |i: &Id| eg[*i].data.constant.clone();
56   match enode {
57     Num(_) | Bool(_) => Some(enode.clone()),
58     Add([x, y]) => Some(Num(c(x)? + c(y)?)),
59     Eq([x, y]) => Some(Bool(c(x)? == c(y)?)),
60     _ => None,
61   }
62 }
63
64 // Functions of this type can be conditions for rewrites
65 trait ConditionFn = Fn(&mut EGraph, Id, &Subst) -> bool;
66
67 // The following two functions return closures of the
68 // correct signature to be used as conditions in Figure 5.1.
69 fn is_not_same_var(v1: Var, v2: Var) -> impl ConditionFn {
70   |eg, _, subst| eg.find(subst[v1]) != eg.find(subst[v2])
71 }
72
73 fn is_const(v: Var) -> impl ConditionFn {
74   |eg, _, subst| eg[subst[v]].data.constant.is_some()
75 }
76
77 struct CaptureAvoid {
78   fresh: Var, v2: Var, e: Var,
79   if_not_free: Pattern<Lambda>, if_free: Pattern<Lambda>,
80 }
81
82 impl Applier<Lambda, LambdaAnalysis> for CaptureAvoid {
83   // Given the egraph, the matching eclass id, and the
84   // substitution generated by the match, apply the rewrite
85   fn apply_one(&self, egraph: &mut EGraph,
86     id: Id, subst: &Subst) -> Vec<Id>
87   {
88     let (v2, e) = (subst[self.v2], subst[self.e]);
89     let v2_free_in_e = egraph[e].data.free.contains(&v2);
90     if v2_free_in_e {
91       let mut subst = subst.clone();
92       // make a fresh symbol using the eclass id
93       let sym = Lambda::Symbol(format!("{}", id).into());
94       subst.insert(self.fresh, egraph.add(sym));
95       // apply the given pattern with the modified subst
96       self.if_free.apply_one(egraph, id, &subst)
97     } else {
98       self.if_not_free.apply_one(egraph, id, &subst)
99     }
100 }

```

Figure 5.2: Our partial evaluator example highlights three important features egg provides for extensibility: e-class analyses, conditional rewrites, and dynamic rewrites.

The `LambdaAnalysis` type, which implements the `Analysis` trait, represents the e-class analysis. Its associated data (`FC`) stores the constant term from that e-class (if any) and an over-approximation of the free variables used by terms in that e-class. The constant term is used to perform constant folding. The merge operation implements the semilattice join, combining the free variable sets and taking a constant if one exists. In `make`, the analysis computes the free variable sets based on the e-node and the free variables of its children; the `eval` generates the new constants if possible. The `modify` hook of `Analysis` adds the constant to the e-graph.

Some of the conditional rewrites in Figure 5.1 depend on conditions defined here. Any function with the correct signature may serve as a condition.

The `CaptureAvoid` type implements the `Applier` trait, allowing it to serve as the right-hand side of a rewrite. `CaptureAvoid` takes two patterns and some pattern variables. It checks the free variable set to determine if a capture-avoiding substitution is required, applying the `if_free` pattern if so and the `if_not_free` pattern otherwise.

5.3 Efficiency

egg’s novel *rebuilding* algorithm (Chapter 3) combined with systems programming best practices makes e-graphs—and the equality saturation use case in particular—more efficient than prior tools.

egg is implemented in Rust, giving the compiler freedom to specialize and inline user-written code. This is especially important as egg’s generic nature leads to tight interaction between library code (e.g., searching for rewrites) and user code (e.g., comparing operators). egg is designed from the ground up to use cache-friendly, flat buffers with minimal indirection for most internal data structures. This is in sharp contrast to traditional representations of e-graphs [Nel80, DNS05] that contains many tree- and linked list-like data structures. egg additionally compiles patterns to be executed by a small virtual machine [dMB07], as opposed to recursively walking the tree-like representation of patterns.

Aside from deferred rebuilding, egg’s equality saturation algorithm leads to implementation-level performance enhancements. Searching for rewrite matches, which is the bulk of running time, can be parallelized thanks to the phase separation. Either the rules or e-classes could be searched in parallel. Furthermore, the once-per-iteration frequency of rebuilding allows egg to establish other performance-enhancing invariants that hold during the read-only search phase. For example, egg sorts e-nodes within each e-class to enable binary search, and also maintains a cache mapping function symbols to e-classes that contain e-nodes with that function symbol.

Many of egg’s extensibility features can also be used to improve performance. As mentioned above, rule scheduling can lead to great performance improvement in the face of “expansive” rules that would otherwise dominate the search space. The Runner interface also supports user hooks that can stop the equality saturation after some arbitrary condition. This can be useful when using equality saturation to prove terms equal; once they are unified, there is no point in continuing. egg’s Runners also support batch simplification, where multiple terms can be added to the initial

e-graph before running equality saturation. If the terms are substantially similar, both rewriting and any e-class analyses will benefit from the e-graph's inherent structural deduplication. The case study in Section 6.2 uses batch simplification to achieve a large speedup with simplifying similar expressions.

Chapter 6

Case Studies

This thesis makes the case that equality saturation (and egg in particular) is the right tool for many program optimization and synthesis tasks. The preceding chapters have articulated the technical novelties, practical features, and quantitative evidence that supports this claim. This chapter provides qualitative evidence through case studies of independently-developed projects¹ from diverse domains that incorporated egg.

Users have built program synthesizers, optimizers, and provers with egg across domains such as floating point accuracy, 3D CAD, deep learning, mutation testing, automatic vectorization, and more. In many cases, the developers had first rolled their own e-graph implementations; egg allowed them to delete code, gain performance, and in some cases dramatically broaden the project’s scope thanks to egg’s speed and flexibility. In addition to gaining performance, all projects use egg’s novel extensibility features like e-class analyses and dynamic rewrites.

¹The case studies feature the following works:

Szalinski [NWA ⁺ 20]	Section 6.1	PLDI 2020	Nandi, Willsey , Anderson, Wilcox, Darulova, Grossman, Tatlock
Herbie [PSSWT15]	Section 6.2	PLDI 2015	Pancheekha, Sanchez-Stern, Wilcox, Tatlock
Tensat [YPW ⁺ 21]	Section 6.3	MLSys 2021	Yang, Phothilimtha, Wang, Willsey , Roy, Pienaar
Ruler [NWZ ⁺ 21]	Section 6.4	Under submission	Nandi, Willsey , Zhu, Saiki, Wang, Anderson, Schulz, Grossman, Tatlock

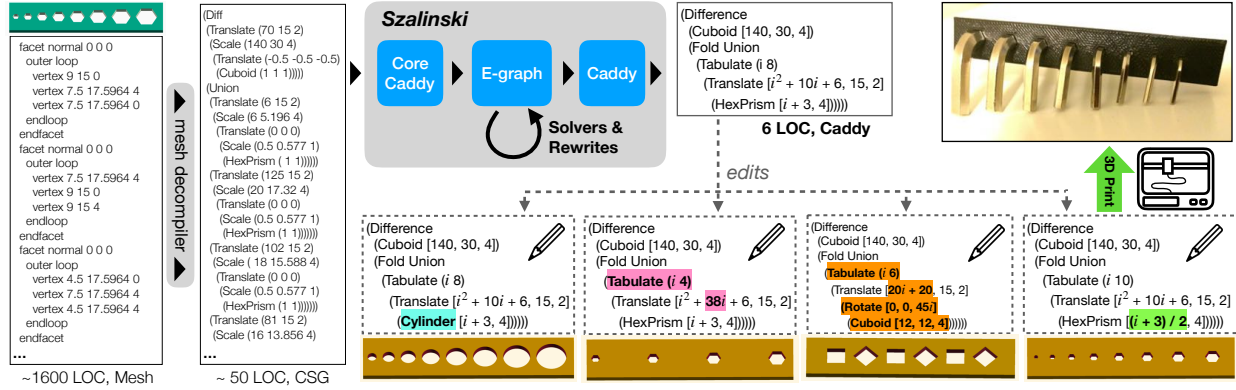


Figure 6.1: (Figure from Nandi et. al. [NWA⁺20]) Existing mesh decompilers turn triangle meshes into flat, computational solid geometry (CSG) expressions. Szalinski [NWA⁺20] takes in these CSG expressions in a format called Core Caddy, and it synthesizes smaller, structured programs in language called Caddy that is enriched with functional-style features. This can ease customization by simplifying edits: small, mostly local changes yield usefully different models. The photo shows the 3D printed hex wrench holder after customizing hole sizes. Szalinski is powered by egg’s extensible equality saturation, relying on its high performance, e-class analyses, and dynamic rewrites.

6.1 Szalinski: Decompiling CAD into Structured Programs

Several tools have emerged that reverse engineer high level Computer Aided Design (CAD) models from polygon meshes and voxels [NWP⁺18, DPIP⁺18, TLS⁺19, SGL⁺17, ERSLT18]. The output of these tools are constructive solid geometry (CSG) programs. A CSG program is comprised of 3D solids like cubes, spheres, cylinders, affine transformations like scale, translate, rotate (which take a 3D vector and a CSG expression as arguments), and binary operators like union, intersection, and difference that combine CSG expressions. For repetitive models like a gear, CSG programs can be too long and therefore difficult to comprehend. A recent tool, Szalinski [NWA⁺20], extracts the inherent structure in the CSG outputs of mesh decompilation tools by automatically inferring maps and folds (Figure 6.1). Szalinski accomplished this using egg’s extensible equality saturation system, allowing it to:

- Discover structure using loop rerolling rules. This allows Szalinski to infer functional patterns like `Fold`, `Map2`, `Repeat` and `Tabulate` from flat CSG inputs
- Identify equivalence among CAD terms that are expressed as different expressions by mesh decompilers. Szalinski accomplishes this by using CAD identities. An example of one such CAD identity in Szalinski is $e \leftrightarrow \text{rotate } [0\ 0\ 0] \ e$. This implies that any CAD expression e is equivalent to a CAD expression that applies a rotation by zero degrees about x, y, and z axes to e
- Use external solvers to speculatively add potentially profitable expressions to the e-graph. Mesh decompilers often generate CSG expressions that order and/or group list elements in non-intuitive ways. To recover structure from such expressions, a tool like Szalinski must be able to reorder and regroup lists that expose any latent structure

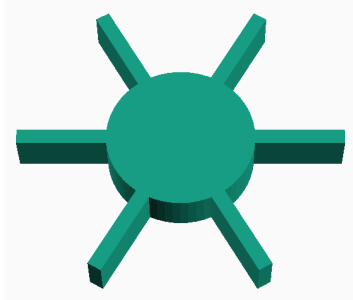
6.1.1 Implementation

Even though CAD is different from traditional languages targeted by programming language techniques, `egg` supports Szalinski’s CAD language in a straightforward manner. Szalinski uses purely syntactic rewrites to express CAD identities and some loop rerolling rules (like inferring a `Fold` from a list of CAD expressions). Critically, however, Szalinski relies on `egg`’s dynamic rewrites and e-class analysis to infer functions for lists.

Consider the flat CSG program in Figure 6.3b. A structure finding rewrite first rewrites the flat list of Unions to:

```
(Fold Union (Map2 Translate [(0 0 0) (2 0 0) ...] (Repeat Cube 5)))
```

The list of vectors is stored as `Cons` elements (sugared above for brevity). Szalinski uses an e-class



(a) CAD model of ship's wheel

```
(Union
  (Cylinder [1, 5, 5])
  (Fold Union
    (Tabulate (i 6)
      (Rotate [0, 0, 60 * i]
        (Translate [1, -0.5, 0]
          (Cuboid [10, 1, 1]))))))))
```

(b) Caddy program

```
(Union
  (Cylinder [1,5])
  (Union
    (Rotate [0,0,0] (Translate [1,-0.5,0] (Cuboid [10,1,1])))
    (Rotate [0,0,60] (Translate [1,-0.5,0] (Cuboid [10,1,1])))
    (Rotate [0,0,120] (Translate [1,-0.5,0] (Cuboid [10,1,1])))
    (Rotate [0,0,180] (Translate [1,-0.5,0] (Cuboid [10,1,1])))
    (Rotate [0,0,240] (Translate [1,-0.5,0] (Cuboid [10,1,1])))
    (Rotate [0,0,300] (Translate [1,-0.5,0] (Cuboid [10,1,1])))))
```

(c) Ideal Core Caddy expression that exposes structure

```
(Union
  (Rotate [0,0,120] (Translate [1,-0.5,0] (Cuboid [10,1,1])))
  (Scale [10,1,1] (Translate [0.1,-0.5,1] (Cuboid [1,1,1])))
  (Rotate [0,0,300] (Translate [1,-0.5,0] (Cuboid [10,1,1])))
  (Scale [5,5,1] (Cylinder [1,1]))
  (Translate [-1,0.5,0] (Scale [-1,-1,1] (Cuboid [10,1,1])))
  (Rotate [0,0,240] (Translate [1,-0.5,0] (Cuboid [10,1,1])))
  (Rotate [0,0,60] (Translate [1,-0.5,0] (Cuboid [10,1,1]))))
```

(d) Equivalent Core Caddy expression that obfuscates structure

Figure 6.2: (a) CAD model for a ship's wheel. (b) Caddy features like Tabulate express repeated design components. Such repetition can be obvious in Core Caddy (c), but existing mesh decompilers obfuscate structure (d).

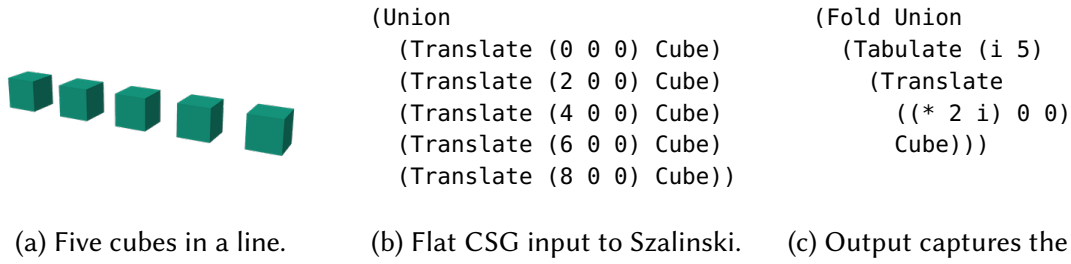


Figure 6.3: Szalinski integrates solvers into egg’s equality saturation as a dynamic rewrite. The solver-backed rewrites can transform repetitive lists into `Tabulate` expressions that capture the repetitive structure.

analysis to track the accumulated lists in a similar style to constant folding. Then, a dynamic rewrite uses an arithmetic solver to rewrite the concrete list of 3D vectors in the analysis data to `(Tabulate (i 5) (* 2 i))`. A final set of syntactic rewrites can hoist the `Tabulate`, yielding the result on the right of Figure 6.3. Thanks to the set of syntactic CAD rewrites, this structure finding even works in the face of CAD identities. For example, the original program may omit the no-op `Translate (0 0 0)`, even though it is necessary to see repetitive structure.

6.1.2 Inverse Transformations

In many cases, the repetitive structure of input CSG expression is further obfuscated because subexpressions may appear in arbitrary order. For these inputs, the arithmetic solvers must first reorder the expressions to find a closed form like a `Tabulate` as shown in Figure 6.3. However, reordering a list does not preserve equivalence, so adding it to the e-class of the concrete list would be unsound. Szalinski therefore introduces *inverse transformations*, a novel technique that allows solvers to speculatively reorder and regroup list elements to find a closed form. The solvers annotate the potentially profitable expression with the permutation or grouping that led to the successful discovery of the closed form. Later in the rewriting process, syntactic rewrites eliminate the inverse transformations when possible (e.g., reordering lists under a `Fold Union`

can be eliminated). `egg` supported this novel technique without modification.

6.1.3 Results

Szalinski’s initial prototype used a custom e-graph written in OCaml. Anecdotally, switching to `egg` removed most of the code, eliminated bugs, facilitated the key contributions of solver-backed rewrites and inverse transformations, and made the tool about 1000× faster. `egg`’s performance allowed a shift from running on small, hand-picked examples to a comprehensive evaluation on over 2000 real-world models from a 3D model sharing forum [NWA⁺20].

6.2 Herbie: Improving Floating Point Accuracy

Herbie automatically improves accuracy for floating-point expressions, using random sampling to measure error, a set of rewrite rules for generating program variants, and algorithms that prune and combine program variants to achieve minimal error. Herbie received PLDI 2015's Distinguished Paper award [PSSWT15] and has been continuously developed since then, sporting hundreds of Github stars, hundreds of downloads, and thousands of users on its online version. Herbie uses e-graphs for algebraic simplification of mathematical expressions, which is especially important for avoiding floating-point errors introduced by cancellation, function inverses, and redundant computation.

Until our case study, Herbie used a custom e-graph implementation written in Racket (Herbie's implementation language) that closely followed traditional e-graph implementations. With timeouts disabled, e-graph-based simplification consumed the vast majority of Herbie's run time. As a fix, Herbie sharply limits the simplification process, placing a size limit on the e-graph itself and a time limit on the whole procedure. When the timeout is exceeded, simplification fails altogether. Furthermore, the Herbie authors knew of several features that they believed would improve Herbie's output but could not be implemented because they required more calls to simplification and would thus introduce unacceptable slowdowns. Taken together, slow simplification reduced Herbie's performance, completeness, and efficacy.

We implemented a egg simplification backend for Herbie. The egg backend is over 3000× faster than Herbie's initial simplifier and is now used by default as of Herbie 1.4. Herbie has also backported some of egg's features like batch simplification and rebuilding to its e-graph implementation (which is still usable, just not the default), demonstrating the portability of egg's conceptual improvements.

6.2.1 Implementation

Herbie is implemented in Racket while egg is in Rust; the egg simplification backend is thus implemented as a Rust library that provides a C-level API for Herbie to access via foreign-function interface (FFI). The Rust library defines the Herbie expression grammar (with named constants, numeric constants, variables, and operations) as well as the e-class analysis necessary to do constant folding. The library is implemented in under 500 lines of Rust.

Herbie’s set of rewrite rules is not fixed; users can select which rewrites to use using command-line flags. Herbie serializes the rewrites to strings, and the egg backend parses and instantiates them on the Rust side.

Herbie separates exact and inexact program constants: exact operations on exact constants (such as the addition of two rational numbers) are evaluated and added to the e-graph, while operations on inexact constants or that yield inexact outputs are not. We thus split numeric constants in the Rust-side grammar between exact rational numbers and inexact constants, which are described by an opaque identifier, and transformed Racket-side expressions into this form before serializing them and passing them to the Rust driver. To evaluate operations on exact constants, we used the constant folding e-class analysis to track the “exact value” of each e-class.² Every time an operation e-node is added to the egg e-graph, we check whether all arguments to that operation have exact value (using the analysis data), and if so do rational number arithmetic to evaluate it. The e-class analysis is cleaner than the corresponding code in Herbie’s implementation, which is a built-in pass over the entire e-graph.

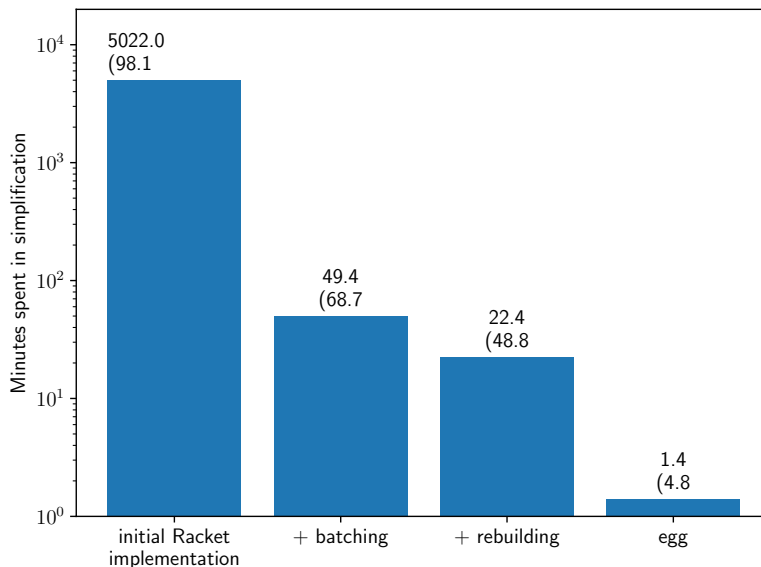


Figure 6.4: Herbie sped up its expression simplification phase by adopting egg-inspired features like batched simplification and rebuilding into its Racket-based e-graph implementation. Herbie also supports using egg itself for additional speedup. Note that the y-axis is log-scale.

6.2.2 Results

Our egg simplification backend is a drop-in replacement to the existing Herbie simplifier, making it easy to compare speed and results. We compare using Herbie’s standard test suite of roughly 500 benchmarks, with timeouts disabled. Figure 6.4 shows the results. The egg simplification backend is over 3000 $\times$ faster than Herbie’s initial simplifier. This speedup eliminated Herbie’s largest bottleneck: the initial implementation dominated Herbie’s total run time at 98.1%, backporting egg improvements into Herbie cuts that to about half the total run time, and egg simplification takes under 5% of the total run time. Practically, the run time of Herbie’s initial implementation was smaller, since timeouts cause tests failures when simplification takes too long. Therefore, the speedup also improved Herbie’s completeness, as simplification now never times out.

²Herbie’s rewrite rules guarantee that different exact values can never become equal; the semilattice join checks this invariant on the Rust side.

Since incorporating `egg` into Herbie, the Herbie developers have backported some of `egg`'s key performance improvements into the Racket e-graph implementation. First, batch simplification gives a large speedup because Herbie simplifies many similar expressions. When done simultaneously in one equality saturation, the e-graph's structural sharing can massively deduplicate work. Second, deferring rebuilding (as discussed in Chapter 3) gives a further $2.2\times$ speedup. As demonstrated in Figure 3.4, rebuilding offers an asymptotic speedup, so Herbie's improved implementation (and the `egg` backend as well) will scale better as the search size grows.

6.3 Tensat: Optimizing Deep Learning Computation Graphs

Deep learning frameworks and compilers (e.g., Tensorflow [ABC⁺16], PyTorch [PGM⁺19], XLA [Goo17], TensorRT [NVI], TVM [CMJ⁺18], MLIR [LAB⁺20]) have enabled diverse kinds of machine learning models to run efficiently on numerous compute platforms. Neural network models in these frameworks are typically represented as tensor computation graphs. To improve the runtime performance of a tensor graph, these frameworks perform various optimizations.

One of the most important optimizations is graph rewriting, which takes in a tensor graph g and a set of semantics-preserving graph rewrites R , and by applying rewrites to g seeks to find an semantically equivalent g' with lower cost according to some cost model. The current industry-standard approach adopted by most frameworks is to use a manually curated set of rewrite rules and rely on a heuristic strategy to determine the order in which to apply the rewrite rules. However, this approach often leads to sub-optimal results both due to the non-comprehensive set of rewrite rules, as well as the sub-optimal graph substitution heuristic [JPT⁺19, JTW⁺19].

This case study aims to address the sub-optimality problem of graph rewrite strategies, while leveraging the existing rewrite rules generation technique [JPT⁺19]. Prior research has shown that searching for sequences of substitutions [JPT⁺19, JTW⁺19, FSWC20] outperforms heuristic approaches. However, both heuristic and search-based solutions rely on sequential application of substitutions. Since rewrites often depend on or enable one another, optimization depends heavily on the order in which rewrites are applied; the “phase ordering” problem strikes again.

This case study presents Tensat, a tensor graph superoptimization framework that employs equality saturation [TSTL09, STL11, WWF⁺20], to apply all possible rewrites at once. Tensat splits program optimization into two phases: *exploration* and *extraction*. The exploration phase is equality saturation using egg as usual. Tensat’s extraction phase is totally custom; simple cost functions simply do not suffice for extracting efficient deep learning compute graphs. Instead,

	Search time (s)		Runtime speedup (%)	
	TASO	Tensat	TASO	Tensat
BERT	13.6	1.4	8.5	9.2
ResNeXt-50	25.3	0.7	5.5	8.8
NasNet-A	1226	10.6	1.9	7.3
NasRNN	177.3	0.5	45.4	68.9
Inception-v3	68.6	5.1	6.3	10.0
SqueezeNet	16.4	0.3	6.7	24.5
VGG-19	8.9	0.4	8.9	8.9

Table 6.1: Comparison of optimization time and runtime speedup of the optimized computation graphs over the original graphs, TASO [JPT⁺19] v.s. Tensat.

Tensat employs an Integer Linear Programming (ILP) extraction solution, which requires a novel method to filter out invalid subgraphs from an e-graph.

We evaluated Tensat on a number of well-known machine learning models executing on a GPU. As highlighted in Table 6.1, Tensat can synthesize optimized graphs that are up to 23% faster in runtime than state-of-the-art [JPT⁺19], while reducing the optimization time by up to 300x. By having the e-graph compactly representing an exponential number of equivalent graphs, Tensat is able to cover a larger search space more efficiently than the sequential search methods. As a result, our search approach is both extremely effective and fast enough to be used as part of a normal compilation flow.

6.3.1 Representation

This section describes how Tensat represents tensor computation graphs and rewrite rules.

Representing Tensor Computation Graphs We use a representation based on the one in TASO [JPT⁺19], with modifications to make it suitable for equality saturation. Table 6.2 shows the set of operators we consider. Each operator o_i corresponds to a node n_i in the graph; the node represents the output tensor of the operator. The nodes corresponding to the inputs of o_i are the

children nodes of n_i . Each tensor computation graph is a DAG under this representation.

The formulations in equality saturation become simpler if a graph is single-rooted. Therefore, we combine all the final output nodes of a graph with *no-ops* to make the graph single-rooted. The no-op nodes do not have any actual operators associated with them, and they will not be altered during the exploration phase, so there are no side effects.

Representing Rewrite Rules A rewrite rule for tensor computation graph specifies that some local subgraph pattern (*source pattern*) is equivalent to another subgraph pattern (*target pattern*). The input tensors to the source and target patterns are *variable nodes*, which can be substituted with any concrete nodes (or e-class in equality saturation) in the current graph. Each output tensor in the source pattern corresponds to an output tensor in the target pattern. The two corresponding output nodes are called a pair of *matched outputs*. A rewrite rule states the equivalence between each pair of matched outputs.

We represent each source (and target) pattern using symbolic expressions (S-exprs) with variables. Patterns with a single output is represented with an S-expr rooted on the output. Rewrite rules with such patterns are called *single-pattern rewrite rules*. Patterns with multiple outputs are represented as a list of S-exprs rooted on each output. Rewrite rules with multiple matched outputs are called *multi-pattern rewrite rules*.

6.3.2 Exploration Phase

We initialize the e-graph with the original tensor computation graph. In each iteration of the exploration phase, we search for matches of all rewrite rules in the current e-graph, and add the target patterns and equivalence relations to the e-graph. This process continues until either the e-graph saturates or a user-specified limit (in terms of time, e-graph size, or number of iterations) is reached. Before applying a rewrite at a found match, we perform a *shape checking* to verify if

the tensor shapes in the target pattern are compatible. This is necessary since some rewrite rules requires input tensor shapes to satisfy specific preconditions, in addition to the syntactic match. We perform shape checking in the same way as TASO [JPT⁺19].

6.3.3 Extraction Phase

During extraction, the goal is to pick one e-node from each e-class in the e-graph to obtain an optimized graph. The optimized graph should minimize the total cost with respect to a given cost model. In tensor graph superoptimization, the cost model reflects the inference time taken by the graph.

Cost model We use the same cost model as TASO [JPT⁺19]. Each operator has a separate and independent cost, which is the measured runtime of that operator (with the specific input sizes and parameters) on hardware. The total cost of a graph is the sum of costs of each of its nodes. This cost model is suitable for GPUs, since GPUs typically run one operator at a time when executing a graph. Note that an operator can be a fused operator, consisting of multiple primitive operators, such as a fused convolution and ReLU.

Greedy extraction We first experiment with a greedy extraction strategy that has been shown to be effective for certain domains [PSSWT15, WHL⁺20, WWF⁺20]. For each e-class, the greedy strategy computes the total cost of the subtrees rooted on each of the e-nodes, and picks the e-node with the smallest subtree cost.

Greedy extraction is not guaranteed to extract the graph with the minimum cost, even under our independent cost model. For example, if two children of an e-node share a subgraph, greedy extraction would ignore the sharing and overestimate the cost.

ILP extraction The second approach we experiment with is formulating the extraction problem as an Integer Linear Program (ILP).

Let $i = 0, \dots, N - 1$ be the set of e-nodes in the e-graph. Let $m = 0, \dots, M - 1$ be the set of e-classes in the e-graph. Let e_m denote the set of e-nodes within e-class m : $\{i | i \in e_m\}$. Let h_i denote the set of children e-classes for e-node i . Let $g(i)$ denote the e-class of e-node i , i.e. $i \in e_{g(i)}$. Let $m = 0$ be the root e-class. Each e-node is associated with a cost c_i .

We then formulate our problem as follows:

$$\text{Minimize: } f(x) = \sum_i c_i x_i$$

Subject to:

$$x_i \in \{0, 1\}, \tag{6.1}$$

$$\sum_{i \in e_0} x_i = 1, \tag{6.2}$$

$$\forall i, \forall m \in h_i, x_i \leq \sum_{j \in e_m} x_j, \tag{6.3}$$

$$\forall i, \forall m \in h_i, t_{g(i)} - t_m - \epsilon + A(1 - x_i) \geq 0, \tag{6.4}$$

$$\forall m, 0 \leq t_m \leq 1, \tag{6.5}$$

Here we introduce a binary integer variable x_i for each e-node i ; node i is selected if $x_i = 1$, and not selected otherwise. Constraint (2) ensures that one node is picked in the root e-class. Constraint (3) ensures that if a node is picked, then at least one node in each of its children e-classes needs to be picked. We rely on the fact that at the optimal solution, each e-class can have at most one picked node (otherwise we can remove more picked nodes in this e-class to reduce the objective while still satisfying all the constraints). Constraints (1)–(3) and the objective encode

the main extraction logic.

A more subtle requirement on the extraction phase is that the extracted graph cannot contain cycles. While the e-graph can (and likely will) contain cycles, the extracted graph is meant to map directly to an executable tensor DAG. The extraction procedure must therefore take care to respect the acyclic invariant of DAGs.

To ensure the extracted graph does not contain cycles, we introduce a real variable t_m for each e-class m in the ILP. Constraint (4) ensures that the order defined by t_m 's is a valid topological order for the extracted graph. Here $\epsilon < 1/M$ is a small constant for effectively encoding strict inequalities in ILP. A is a large enough constant such that $A > 1 + \epsilon$. Constraint (5) is to limit the range for the topological order variables t_m 's.

We also experiment with using integer variables for t_m 's. In this case, t_m 's are constrained to take integer values between 0 to $M - 1$. Constraint (4) changes accordingly to: $\forall i, \forall m \in h_i, t_{g(i)} - t_m + A(1 - x_i) \geq 1$, where $A \geq M$.

Unlike greedy extraction, the optimal solution to the ILP is guaranteed to give a valid graph (no cycles) with the lowest cost.

Cycle Filtering Similar to previous work that uses ILP extraction [TSTL09, WHL⁺20], we find that as the size of the e-graph grows bigger, the ILP solver takes a long time and becomes the main bottleneck. This is mainly due to the cycle constraint (4): ILP solver struggles to find a feasible solution with these constraints. Therefore, we explore an alternative approach by filtering cycles during the exploration phase to make sure that the e-graph does not contain any cycles at the end of the exploration phase. This way, we can get rid of the cycle constraints in the ILP.

Vanilla cycle filtering The first method is to check if applying a substitution introduces cycles to the e-graph, and discard such a substitution. This check is run every time before applying a

substitution. Each check requires a pass over the entire e-graph. For one iteration during the exploration phase, if we denote N as the current size of the e-graph and n_m as the total number of matches of the rewrite rules on the e-graph, then this vanilla cycle filtering has complexity $O(n_m N)$.

Efficient cycle filtering As the number of matches n_m is typically large and scales with N , vanilla cycle filtering can be slow. We therefore design a novel and more efficient cycle filtering algorithm, consisting of a *pre-filtering* step and a *post-processing* step. Algorithm 1 shows the pseudocode for the exploration phase with efficient cycle filtering.

At the start of each iteration, we do one pass over the e-graph to record the set of descendent e-classes for each e-node (stored in a descendants map). During the iteration, for each match of the rewrite rules, we use the pre-stored descendants map to check if applying a rewrite introduces cycles to the e-graph; if so, we skip this match. Line 3–9 implements the pre-filtering step. Notice that this check is sound but not complete: a match that passes this check can still introduce cycles to the e-graph. This is because new descendants relations introduced by the previous rewrite in this iteration are not included in the pre-stored descendants map.

To resolve the cycles we missed in the pre-filtering step, we add a post-processing step at the end of each iteration (line 10-18). We make a pass over the e-graph in DFS order and collect a set of cycles in the e-graph. For each cycle, we choose the last node that is added to the e-graph, and add that node to a filter list. The nodes in the filter list are considered as removed from the e-graph. We make sure those nodes are not picked during extraction by explicitly adding constraints $\forall i \in l, x_i = 0$ to the ILP.

By constructing a descendants map once before each iteration, each of the checking in the pre-filtering step takes constant time. The worst case complexity of the post-processing step is $O(n_c N)$, where n_c is the number of cycles in the e-graph. Since n_c is typically much smaller than

Algorithm 1 Exploration phase with efficient cycle filtering

Input: starting e-graph $\mathcal{G}$, set of rewrite rules $\mathcal{R}$.

Output: updated e-graph $\mathcal{G}$, filter list l

```

1:  $l = \{\}$ 
2: for iter = 0, ..., MAX_ITER do
3:   descendants map  $d = \text{GETDESCENDANTS}(\mathcal{G}, l)$ 
4:   matches = SEARCH( $\mathcal{G}, \mathcal{R}, l$ )
5:   for match  $\in$  matches do
6:     if not WILLCREATECYCLE(match,  $d$ ) then
7:       APPLY( $\mathcal{G}$ , match)
8:     end if
9:   end for
10:  while true do
11:    cycles = DFSGETCYCLES( $\mathcal{G}, l$ )
12:    if len(cycles) == 0 then
13:      break
14:    end if
15:    for cycle  $\in$  cycles do
16:      RESOLVECYCLE( $\mathcal{G}, l$ , cycle)
17:    end for
18:  end while
19: end for
20: return  $\mathcal{G}, l$ 

```

n_m , this algorithm is much faster than the vanilla cycle filtering. In practice, each DFS pass over the e-graph can find many cycles, which makes $\mathcal{O}(n_c N)$ a very conservative upper bound.

6.3.4 Evaluation

We implemented Tensat in Rust [Rus] using egg [WWF⁺20]. For the extraction phase, we use SCIP [GAB⁺20] as the ILP solver, wrapped by Google OR-tools [PF].

We utilize egg's *e-class analysis* feature for the shape checking discussed. An e-class analysis associates data with each e-class to support rewrites that are not purely syntactic. We store all the relevant information of the tensors (shape, layout, split locations) in the analysis data and use these information for shape checking.

Experimental Setup We compared Tensat with TASO [JPT⁺19] to evaluate our equality saturation based search. We used the same set of rewrite rules as TASO for our experiments. We evaluated on the inference graphs of 7 models: **BERT** [DCLT19], **ResNeXt-50** [XGD⁺17], **NasNet-A** [ZVSL18], **NasRNN** [ZL17], **Inception-v3** [SVI⁺16], **VGG-19** [LD15], and **SqueezeNet** [IMA⁺17]. This benchmark set covers a wide range of commonly used state-of-the-art models, including both models for computer vision tasks and models for NLP tasks, both human-designed models and automatically-discovered models by neural architecture search. We performed all experiments on a Google Cloud instance with one NVIDIA Tesla T4 GPU, a 16-core CPU, and 60 GB of memory.

For Tensat, our full approach uses the efficient cycle filtering algorithm (Section 6.3.3) during the exploration phase and the ILP method without the cycle constraints (Section 6.3.3) for extraction. We set a limit on the number of nodes in the e-graph $N_{\max} = 50000$ and the number of iterations for exploration $k_{\max} = 15$. We terminate the exploration phase when any of the limit is reached, or the e-graph is saturated. We set a separate limit k_{multi} on the number of iterations to apply the multi-pattern rules. We use a default of $k_{\text{multi}} = 1$ for the main results in Section 6.3.4 and Section 6.3.4. We set a timeout of 1 hour for the ILP solver.

For TASO’s backtracking search, we use their default settings from their artifact evaluation code on the number of iterations³ $n = 100$ and the hyperparameter $\alpha = 1.0$ for each benchmark. We also test $\alpha = 1.05$ as mentioned in their paper, and find that the difference is tiny (difference in speedup percentage is less than 0.1% on average over the benchmarks). Increasing to $n = 1000$ leads to less than 1% speedup gain with the cost of over 11x longer in optimization time on average.

Program Speedup We compare the speedup percentage of the optimized graph with respect to the original graph between Tensat and TASO. We use TASO’s cuDNN backend to measure the runtime of the full computation graphs. Figure 6.5 shows the results. We can see that Tensat

³The number of iterations of the outer loop, see Algorithm 2 in [JPT⁺19] for more details

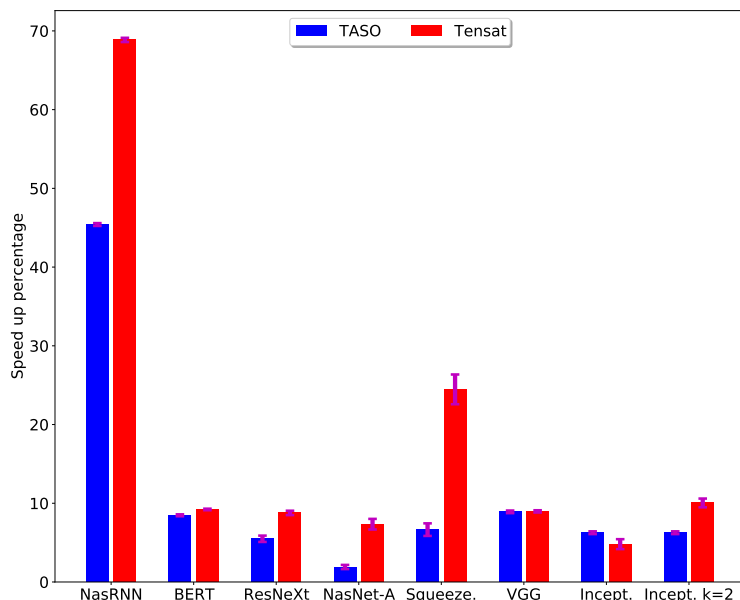


Figure 6.5: Speedup percentage of the optimized graph with respect to the original graph, TASO v.s. Tensat. Each setting (optimizer $\times$ benchmark) is run for five times, and we plot the mean and standard error for the measurements.

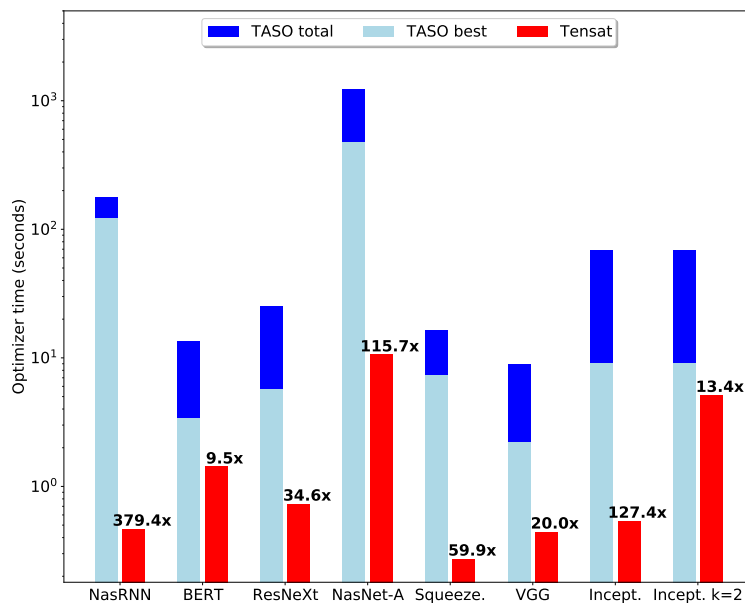


Figure 6.6: Comparison of the optimization time (log scale) between TASO and Tensat. “TASO total” is the total time of TASO search. “TASO best” indicates when TASO found its best result; achieving this time would require an oracle telling it when to stop.

discovers better optimized graphs compared with TASO’s backtracking search in most benchmarks. Tensat’s optimized graphs are on average 6.6% faster than TASO’s. We see the biggest speedup of 23% over TASO on NasRNN. Note that for Inception-v3, Tensat with $k_{\text{multi}} = 1$ gives a smaller speedup than TASO, but increasing k_{multi} to 2 achieves a better speedup than TASO while still being 13.4× faster than TASO’s search (see Figure 6.6).

This improvement comes from the fact that equality saturation covers a much larger space of equivalent graphs than sequential backtracking search. By using e-graph as a compact representation of an exponential number of equivalent graphs, Tensat is able to cover orders of magnitude more equivalent graphs than TASO.

Optimization Time Another important metric is the time taken by the optimizer itself. For Tensat, this is the sum of time taken by the exploration phase and the extraction phase. For TASO, we record two times for a single backtracking search. The first is the total time of the backtracking search with the default number of iterations (T_{total}). The second one is the time taken to first reach the best graph found during its search (T_{best}). T_{best} is the best possible time for TASO’s sequential backtracking search. In practice, it is difficult (if not impossible) to achieve T_{best} since the sequential search algorithm would have no way to know that it can stop at that point.

Figure 6.6 shows the time taken by the optimizers across benchmarks. We can see that Tensat runs 9.5x to 379x faster than TASO’s T_{total} , and 1.8x to 260x times faster than T_{best} . This shows that Tensat can not only cover a much larger search space, but also achieve this in drastically less time. Furthermore, Tensat’s optimization time is small enough that we believe our approach can be integrated into a default compilation flow instead of running the search as an additional offline autotuning process.

Types

T	Tensor	TT	Tensor tuple	W	Weights tensor	s_h, s_w	stride (height, width)
n	natural number	p	padding	a	activation	k_h, k_w	kernel (height, width)

Operator	Description	Type signature
ewadd	Element-wise addition	$(T, T) \rightarrow T$
ewmul	Element-wise multiplication	$(T, T) \rightarrow T$
matmul	Matrix multiplication	$(a, T, T) \rightarrow T$
conv ^a	Grouped convolution	$(s_h, s_w, p, a, T, W) \rightarrow T$
relu	Relu activation	$T \rightarrow T$
tanh	Tanh activation	$T \rightarrow T$
sigmoid	Sigmoid activation	$T \rightarrow T$
poolmax	Max pooling	$(T, k_h, k_w, s_h, s_w, p, a) \rightarrow T$
poolavg	Average pooling	$(T, k_h, k_w, s_h, s_w, p, a) \rightarrow T$
transpose ^b	Transpose	$(T, \text{permutation}) \rightarrow T$
enlarge ^c	Pad a convolution kernel with zeros	$(T, T_{\text{ref}}) \rightarrow T$
concat _{n}	Concatenate along the given axis	$(n, T, \dots, T) \rightarrow T$
split ^d	Split a tensor into two along the axis	$(n, T) \rightarrow TT$
split ₀	Get the first output from split	$TT \rightarrow T$
split ₁	Get the second output from split	$TT \rightarrow T$
merge ^e	Update weight to merge grouped conv	$(W, n) \rightarrow W$
reshape	Reshape tensor	$(T, \text{shape}) \rightarrow T$
input	Input tensor	identifier $\rightarrow T$
weight	Weight tensor	identifier $\rightarrow T$
no-op	Combine the outputs of the graph	$(T, T) \rightarrow T$

Table 6.2: Operators supported by Tensat. There are four types for the nodes in our representation: tensor type (T), string type (S), integer type (N), and tensor tuple type (TT). The integer type is used to represent parameters of the operators, such as stride, axis, and also padding and activation modes (by representing different modes using different integers). The more complex, variable-length parameters (e.g. shape, axes permutation) are represented using the string type according to the specified formats.

^a Same representation as TASO [JPT⁺19]. Normal and depth-wise convolutions are special cases of grouped convolutions.

^b Axis permutation for transpose is specified using a string with format: axis₁_axis₂_...

^c Pad a convolution kernel (input) with zeros to make it the same size as input T_{ref}

^d Split the tensor in the given axis. The position of the split is at the place of the most recent concat.

^e Merge every *count* number of groups in the grouped convolution. See TASO [JPT⁺19] for more details.

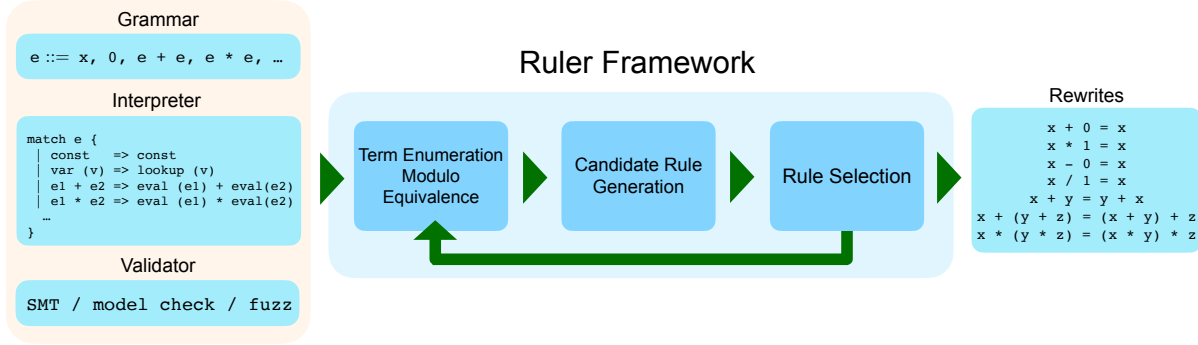


Figure 6.7: *Ruler Workflow*. Given a grammar and interpreter for a target domain, Ruler uses e-graphs and equality saturation to efficiently enumerate potential rewrite rules and iteratively select a small set of general, orthogonal rules. Ruler supports various validation strategies to ensure soundness and speed up synthesis, including constraint solving (e.g., SMT), model checking, and fuzzing.

6.4 Ruler: Rewrite Synthesis using Equality Saturation

Many compilers, program synthesizers, and theorem provers rely on rewrite systems [PJTH01, HA00, DNS05]. For example, rewriting is essential for improving program analyses and code generation [Bli13, LAB⁺20, RKBA⁺13, CMJ⁺18] and for automating verification [BCD⁺11, DMB08, NWP02, BC10]. Without rule-based simplification, Halide-generated code can suffer 26 $\times$ slowdown [NJK⁺20] and the Herbie floating-point synthesizer [PSSWT15] can return 10 $\times$ larger programs.

Where do the rewrite rules come from? Several noteworthy projects have developed tool-specific techniques for checking or inferring rules [BA06, MN17, JNR02, SSL16], but implementing a rewrite system still generally requires domain experts to first manually develop rulesets by trial and error. Such slow, ad hoc, and error-prone approaches hinder design space exploration for new domains and discourage updating existing systems.

To address these challenges, we propose a simple, domain-general approach that uses equality saturation [TSTL09, WWF⁺20] as a rewrite system *on the domain of rewrite rules themselves* to

quickly synthesize effective rulesets.

In the past, tool-specific techniques to iteratively infer rewrite rules have implicitly adopted a common three-step approach, each constructing or maintaining a set:

1. Enumerate terms from the given domain to build the *term set* T .
2. Select candidate rules from $T \times T$ to build the *candidate set* C .
3. Filter C to select a sound set of useful rules to build the *rule set* R .

We identify and abstract this workflow to provide generic rule inference for user-specified domains.

Our key insight is that *what makes equality saturation successful in rewrite rule application is also useful for rule inference*. Equality saturation can simultaneously prove many pairs of terms equivalent with respect to a given ruleset. Ruler uses equality saturation to shrink the set T of enumerated terms (lowering candidate *generation* cost) by merging terms equivalent under R , and to shrink the set C of candidate rules (lowering candidate *selection* cost) by removing rules derivable by R . Thus, Ruler uses the set R of rewrite rules to rewrite the next batch of candidate rewrite rules *even as R is being synthesized*.

We prototyped these insights in a tool dubbed Ruler (Figure 6.7). Compared to a state-of-the-art rule synthesizer [NRB⁺19] built into the CVC4 theorem prover [BCD⁺11], Ruler synthesizes smaller rulesets in less time without reducing the set of derivable equivalences. We demonstrate how Ruler can generate expert-quality rulesets by using it to replace all of Herbie’s rules for rational numbers, uncovering missing rules that resolved a known bug in Herbie.

This case study’s contributions include:

- A novel rule synthesis algorithm that uses e-graphs [Nel80] to compactly encode large sets of terms and equality saturation to efficiently filter and minimize rulesets (Section 6.4.1).

- A generic implementation of this algorithm within the Ruler rewrite rule inference framework that synthesizes rules for user-specified domains given a grammar and its interpreter.
- A comparison against a recent CVC4-based rule synthesizer that shows Ruler synthesizes $5.8\times$ smaller rulesets $25\times$ faster without compromising the deriving power of the rulesets.
- A case study demonstrating that, in an end-to-end application of a real world tool, Ruler’s automatically generated rulesets are as good as manually-crafted expert rules (Section 6.4.4).

We implemented Ruler in Rust using `egg` [WWF⁺20] for equality saturation. `egg`’s flexibility allows Ruler to be relatively simple: its core consists of under 1,000 lines of code, allowing it to be simple, extensible, and generic over domains. Compared to the rewrite synthesis tool inside the CVC4 solver [BCD⁺11, NRB⁺19], Ruler is an order of magnitude smaller.

6.4.1 Ruler’s Algorithm

Like other rule synthesis approaches, Ruler iteratively performs three steps:

1. Enumerate terms into a set T .
2. Search $T \times T$ for a set of candidate equalities C .
3. Choose a useful, valid subset of C to add to the ruleset R .

Ruler’s core insight is that e-graphs and equality saturation can help compactly represent the sets T , C , and R , leading to a faster synthesis procedure that produces smaller rulesets R with greater proving power.

Figure 6.8 shows Ruler’s core synthesis algorithm, which is parameterized by the following:

- The number of iterations to perform the search for (line 4);

```

1 def ruler (iterations):
2      $T$  = empty_egraph()
3      $R$  = {}
4     for  $i \in [0, \text{iterations}]$ :
5         # add new terms directly to the e-graph representing  $T$ 
6         add_terms( $T$ ,  $i$ )
7         loop:
8             # combine e-classes in the e-graph representing  $T$  that  $R$  proves equivalent
9             run_rewrites( $T$ ,  $R$ )
10             $C$  = cvec_match( $T$ )
11            if  $C = \{\}$ :
12                break
13            # choose_eqs only returns valid candidates by using 'is_valid' internally
14             $R = R \cup \text{choose_eqs}(R, C)$ 
15    return  $R$ 

```

Figure 6.8: Ruler’s Core Algorithm. The iterations parameter determines the maximum number of connectives in the terms Ruler will enumerate.

- The language grammar, given in the form of a term enumerator (add_terms, line 6), which takes the number of variables or constants to enumerate over;
- The procedure for validating candidate rules, is_valid (called inside choose_eqs, Figure 6.9 line 20).

These parameters provide flexibility for supporting different domains, making Ruler a rule synthesis framework rather than a single one-size-fits-all tool.

Ruler uses an e-graph to compactly represent the set of terms T . In each iteration, Ruler first extends the set T with additional terms from the target language. Each term $t \in T$ is tagged with a *characteristic vector* (cvec) that stores the result of evaluating t given many different assignments of values to variables.

After enumerating terms, Ruler uses equality saturation (run_rewrites) to merge terms in T that can be proved equivalent by the rewrite rules already discovered (in the set R);

Next, Ruler computes a set C of candidate rules (`cvec_match`). It finds pairs $(t_1, t_2) \in T \times T$ where t_1 and t_2 are from distinct e-classes but have matching cvecs and thus are likely to be equivalent. Thanks to `run_rewrites`, no candidate in C should be derivable from R . However, C is often still large and contains many redundant or invalid candidate rules.

Finally, Ruler's `choose_eqs` procedure picks a valid subset of C to add to R , ideally finding the smallest extension which can establish all equivalences implied by $R \cup C$. Ruler tests candidate rules for validity using a domain-specific `is_valid` function. This process is repeated until there are no more equivalences to learn between terms in T , at which point Ruler begins another iteration.

6.4.2 Choosing Rules

After finding a set of candidate rules C , Ruler selects a valid subset of rules from C to add to the rule set R using the `choose_eqs` procedure (Figure 6.8, line 14). As long as `choose_eqs` returns a valid, non-empty subset of C , Ruler's inner loop will terminate: the number of e-classes with matching cvecs (i.e., the subset of T used to compute C) decreases in each iteration since R is repeatedly extended with rules that will cause new merges in `run_rewrites`. Ideally, `choose_eqs` quickly finds a minimal extension of R that enables deriving all equivalences implied by $R \cup C'$ where C' is the valid subset of C .

The candidate rules in C are not derivable by R , but many of the candidate rules may be able to derive each other, especially in the context of R . For example, the following candidate set is composed of three rules from the boolean domain, and any two can derive the third:

$$(\wedge x x) = \text{false} \qquad (\& x \text{false}) = \text{false} \qquad (\& x \text{false}) = (\wedge x x)$$

An implementation of `choose_eqs` that only returns a single rule $c \in C$ avoids this issue, since adding c to R prevents those rules derivable by $R \cup \{c\}$ from being candidates in the next iteration of the inner loop. However, a single-rule implementation will be slow to learn rules, since it can

```

1  # R is the accepted ruleset so far, C is the candidate ruleset.
2  # Ruler's implementation of choose_eqs is based on a more flexible choose_eqs_n.
3  def choose_eqs(R, C, n = ∞):
4      for step ∈ [100, 10, 1]:
5          if step ≤ n:
6              C = choose_eqs_n(R, C, n, step)
7      return C
8
9  # n is the number of rules to choose from C, and step is a granularity parameter.
10 # A larger step size allows you to eliminate redundant rules faster.
11 def choose_eqs_n(R, C, n, step):
12     # let K be the list of "keepers" which we will return
13     K = []
14     while C ≠ ∅:
15         # pick the best step candidate rules from C according to a heuristic
16         # that approximates rule "generality", including subsumption.
17         Cbest, C = select(step, C)
18
19         # add the valid ones to K
20         K = K ∪ {c | c ∈ Cbest. is_valid(c)}
21
22         # remember all the invalid candidates in a global variable bad;
23         # Ruler uses this to prevent known-invalid candidates from entering C again (not shown)
24         bad = bad ∪ {c | c ∈ Cbest. ¬is_valid(c)}
25
26         # stop if we have enough rules
27         if |K| ≥ n:
28             return K[0..n]
29
30         # try to prove terms remaining in C equivalent using rules from R ∪ K
31         C = shrink(R ∪ K, C)
32     return K
33
34 def shrink(R, C):
35     E = empty_egrph()
36     for (l → r) ∈ C:
37         E = add_term(E, l)
38         E = add_term(E, r)
39     E = run_rewrites(E, R)
40     # return the extracted versions of rules from C, leaving out anything that was proven equivalent
41     return {extract(E, l) → extract(E, r) | (l → r) ∈ C. ¬equiv(E, l, r)}

```

Figure 6.9: Ruler's implementation of choose_eqs, which aims to minimize the candidate set C by eliminating subsets that the remainder can derive.

only learn one at a time (Table 6.3 of our evaluation shows there are sometimes thousands of rules to learn). Additionally, such an implementation has to decide which rule to select, ideally picking the “strongest” rules first. For example, if $a, b \in C$ and $R \cup \{a\}$ can derive b but $R \cup \{b\}$ can not derive a , then selecting b before a would be a mistake, causing the algorithm to incur an additional loop.

Ruler’s implementation of `choose_eqs`, shown in Figure 6.9, is parameterized by a value n with default of ∞ . At $n = 1$, `choose_eqs` simply returns a single valid candidate from C . For higher n , `choose_eqs` attempts to return a list of up to n valid rules all at once. This can speed up Ruler by requiring fewer trips around its inner loop, but risks returning many rules that can derive each other. To mitigate this, `choose_eqs` tries to not choose rules that can derive each other. In its main loop (line 14), `choose_eqs` uses the `select` function to pick the *step* best rules from C according to a syntactic heuristic.⁴ Ruler then validates the selected rules and adds them to a set K of “keeper” rules which it will ultimately return. It then employs the `shrink` procedure (line 34) to eliminate candidates from C that can be derived by $R \cup K$. This works similarly to `run_rewrites` in the Ruler algorithm, but `shrink` works over the remaining *candidate set* C instead of the rule set R .

Ruler’s `choose_eqs` invokes the inner `choose_eqs_n` procedure with increasing small step sizes (*step* is defined on line 4). Larger step sizes allow `shrink` to quickly “trim down” C when it contains many candidates. However, a large step also means that `choose_eqs` may admit *step* rules into K at once, some of which may be able to prove each other. Decreasing the step size to 1 eliminates this issue.

⁴Ruler’s syntactic heuristic prefers candidates with the following characteristics (lexicographically): more distinct variables, fewer constants, shorter larger side (between the two terms forming the candidate), shorter smaller side, and fewer distinct operators.

6.4.3 Comparison with CVC4

To evaluate Ruler, we compared it with prior work that synthesizes rewrites using the CVC4 solver [NRB⁺19]. Both Ruler and the CVC4 synthesizer are written in systems programming languages (Rust and C++, respectively), and both take similar approach to synthesizing rewrite rules: enumerate terms, find valid candidates, select rules and repeat.

We compared Ruler against CVC4 for booleans, bitvector-4, and bitvector-32. Both Ruler and CVC4 are parameterized by the domain (bool, bv4, or bv32), the number of distinct variables in the grammar, and the size of the synthesized term.⁵ All benchmarks were single-threaded and run on an AMD 3900X 3.6GHz processor with 32GB of RAM. Both Ruler and CVC4 were given 3 variables and no constants to start the enumeration.

A bigger ruleset is not necessarily a better ruleset. We designed Ruler to minimize ruleset size while not compromising on its capability to prove equalities. We define a metric called the *deriving ratio* to compare two rulesets. Ruleset A has deriving ratio p with respect to ruleset B if set A can derive a fraction p of the rules in B ($A \models b$ means rule set A can prove rule b):

$$p = |B_A|/|B| \quad \text{where} \quad B_A = \{b \mid b \in B. A \models b\}$$

If A and B have deriving ratio of 1 with respect to each other, then they can each derive all of the other's rules.

We use egg's equality saturation procedure to test derivability. To test whether $A \models b$ (where $b = b_l \rightarrow b_r$) we add b_l and b_r to an empty e-graph, run equality saturation using A , and check to see if the e-classes of b_l and b_r merged. We run egg with 5 iterations of equality saturation. Since this style of proof is bidirectional (egg is trying to rewrite both sides at the same time), derivations

⁵Size is measured in number of connectives, e.g., a has 0, $(a + b)$ has 1, and $(a + (b + c))$ has 2. In CVC4, this is set with the `-sygus-abort-size` flag.

Parameters		Ruler			CVC4			Ruler / CVC4	
Domain	# Conn	Time (s)	# Rules	Drv	Time (s)	# Rules	Drv	Time	Rules
bool	2	0.01	20	1	0.13	53	1	0.06	0.38
bool	3	0.06	28	1	0.82	293	1	0.07	0.10
bv4	2	0.14	49	1	4.47	135	0.98	0.03	0.36
bv4	3	4.30	272	1	372.26	1978	1	0.01	0.14
bv32	2	13.00	46	0.97	18.53	126	0.93	0.70	0.37
bv32	3	630.09	188	0.98	1199.53	1782	0.91	0.53	0.11
								0.04	0.17
								Harmonic Mean	

Table 6.3: Ruler tends to synthesize smaller, more powerful rulesets in less time than CVC4. The table shows synthesis results across domains, and number of variables in the grammar, and maximum term size (in number of connectives, “# Conn”). The domains are booleans, bitvector-4, and bitvector-32. For verification, Ruler uses model checking for booleans and bitvector-4 and Z3 for bitvector-32. The “Drv” column shows the fraction that tool’s synthesized ruleset can derive of the other’s ruleset; for example, the final row indicates that Ruler’s 188 rules derived 98% of CVC4’s 1,782 rules, and CVC’s rules derived 91% of Ruler’s. The final two columns show the ratios of synthesis times and ruleset sizes between the two tools.

of $b_l = b_r$ can be as long as 10 rules from A .

Table 6.3 shows the results of our comparison with CVC4’s rewrite rule synthesis. On average (harmonic mean), Ruler produces $5.8\times$ smaller rulesets $25\times$ faster than CVC4. Ruler and CVC4’s results can derive each most of other. On the harder benchmarks (in terms of synthesis times), Ruler’s results have a higher derivability ratio; they can prove more of CVC4 rules than vice-versa.

6.4.4 Synthesizing Herbie Rewrites

We also demonstrate that Ruler-generated rules can replace and augment those generated by experts by doing exactly that for the Herbie tool [PSSWT15] (described in Section 6.2).

Experimental Setup We implemented rational numbers in Ruler, synthesized rewrite rules over rational arithmetic, and then ran Herbie with the resulting ruleset.

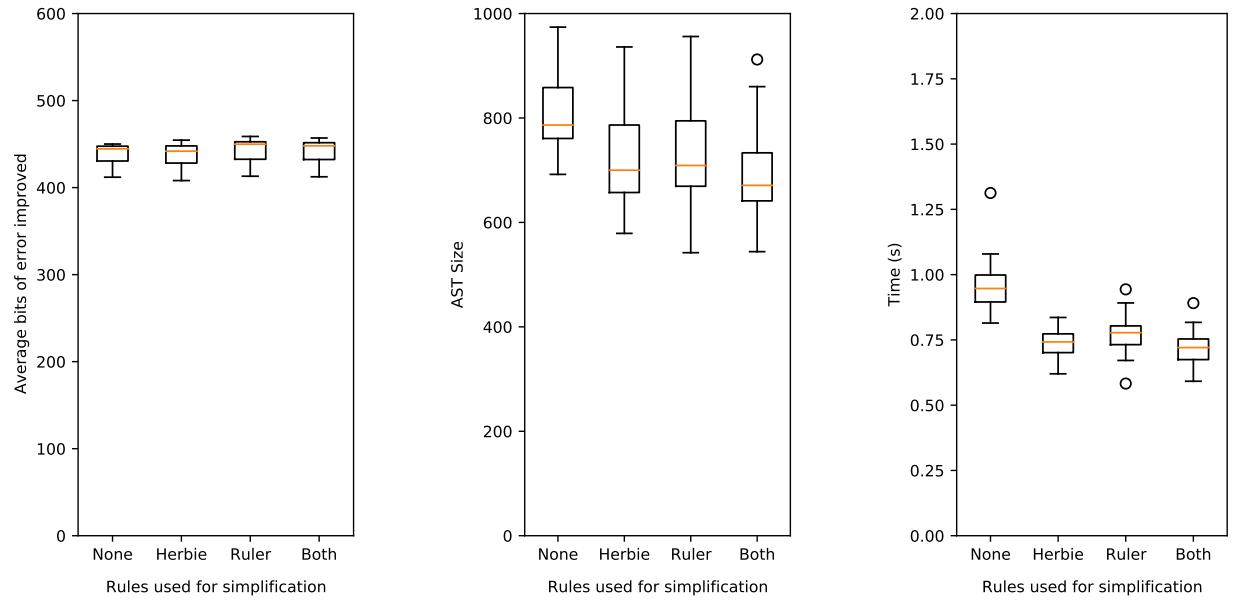
The Herbie benchmark suite has 51 stable benchmarks that contain only rational operators (as opposed to things like $\sin$ and $\cos$). We ran Herbie on these benchmarks under four different configurations:

- **None:** remove all the rational rewrite rules from Herbie’s simplification phase. Rational rules are those which consist only of rational operators and no others. Note that all other components of Herbie are left intact, including rules over rational operators combined with other operators, and rules entirely over other operators. None is the baseline.
- **Herbie:** no changes to Herbie, simply run it on the 51 benchmarks.
- **Ruler:** replace Herbie’s rational rules with output of Ruler.
- **Both:** run Herbie with both Ruler’s rational rules and the original Herbie rational rules.

We used Ruler to synthesize rational rules of depth 2 with 3 variables.⁶ Ruler learned 50 rules in 18 seconds, all of which were proven sound with an SMT post-pass. Four rules were expansive — i.e., rules like $(a \rightarrow (a \times 1))$ whose LHS is only a variable. We removed these expansive rules from the ruleset as per the recommendation of the Herbie developers.

Discussion The Herbie simplifier uses equality saturation to find smaller, equivalent programs. The simplifier itself does not directly improve accuracy; rather, it generates more candidates that are then used in the other accuracy improving components of Herbie. While ideally, Herbie would return a more accurate *and* smaller output, Herbie’s ultimate goal is to find more accurate expressions, even if it sacrifices AST size. Herbie’s original ruleset has been developed over the past 6 years by numerical methods experts to effectively accomplish this goal. Any change to these rules must therefore ensure that it does not make Herbie’s result less accurate.

⁶For rationals, the `add_terms` implementation enumerates terms by depth rather than number connectives, since that matches the structure of Herbie’s existing rules.



(a) Improvement in average error, Herbie's metric for measuring accuracy (higher is better).

(b) Size of the output AST produced by Herbie (lower is better).

(c) Herbie's running time (lower is better).

Figure 6.10: Comparing Herbie results between four configurations. Each boxplot represents the results from 30 seeds, where each data point is obtained by summing the value (average error, AST size, time) over all 51 benchmarks. The columns dictate what rational rules Herbie has access to: either none, its default rules, only Ruler's rules, or both. Herbie's rational rules reduce AST size and speed up simplification without reducing accuracy, and Ruler's rules perform similarly (with or without Herbie's rules).

Figure 6.10 shows the results of running Herbie with rules synthesized by Ruler. Each box-plot corresponds to one of the four configurations. The baseline (None) and Herbie in Figure 6.10's accuracy and AST size plots highlight the significance of rational rewrites in Herbie — these expert-written rules reduce AST size without reducing accuracy. The plots for Ruler show that running Herbie with only Ruler's rational rules has almost the same effect on accuracy and AST size as Herbie's original, expert written ruleset. The plot for Both shows that running Herbie together with Ruler's rules further reduces AST size, still without affecting accuracy. The timing plots show that adding Ruler's rules to Herbie does not make it slower. The baseline timing is slower than the rest because removing all rational simplification rules causes Herbie's other components take much longer to find the same results.

In summary, Ruler's rational rewrite rules can be easily integrated into Herbie, and they perform as well as expert-written rules without incurring any additional overhead.

Fixing a Herbie Bug Ruler found the following two rules that helped the Herbie team address a GitHub issue [Her21]: $(|a \times b| \rightarrow |a| \times |b|)$, and $(|a \times a| \rightarrow a \times a)$. In many cases, Herbie may generate large, complex outputs without improving accuracy, which makes the program unreadable and hard to debug. This is often due to lack of appropriate rules for expression simplification. The issue raised by a user ([Her21]) was in fact due to the missing rule $(|x| \times |x| \rightarrow x \times x)$. The two rules above, can together, accomplish the effect of this rule, thereby solving the issue. We submitted these two rules to the Herbie developers and they added them to their ruleset.

Chapter 7

Conclusion

To conclude this thesis, we should return to our original goal of supporting the thesis statement stated in Chapter 1:

E-graphs and equality saturation are compelling techniques for program representation and manipulation that should now be considered for programming tools across many domains.

There are a few important parts of this statement, and teasing them apart may help us figure out if we have done what we set out to do.

First, we start out with “e-graphs and equality saturation *are* compelling techniques [...]”. This introductory phrase deserves to stand alone. Both the core technique discussed in this thesis and the data structure that powers it are exciting, well-developed prior work. My hope is that the document as a whole—and the background in Chapter 2 in particular—get this point across to the reader. In the course of advancing the state-of-the-art in these areas, it is necessary to focus on their shortcomings, but hopefully this phrase expresses how just how tall these shoulders are. Had I not been excited by these works, this thesis would not exist.

Second, we have “e-graphs and equality saturation [...] should *now* be considered [...]”. This is not to say that they are not without merit on their own, but rather that I hope the advances introduced in this thesis help raise the profile of equality saturation as a compelling technique. Chapters 3 and 4 present new approaches that makes equality saturation faster and more flexible, alleviating two key concerns with that a prospective user may have. Chapter 5 introduces *egg*, the tool that implements all of this and make equality saturation easier to use than ever before. Put together, I believe these contributions make *now* the time to look into and use equality saturation.

Finally, we end by stating that equality saturation should be applicable to “[...] programming tools across many domains.” My goal for this thesis is to turn many “why” questions about equality saturation into “why not” questions. The case studies in Chapter 6 hopefully offer empirical evidence that equality saturation can be useful (and even critical) in unexpected ways. If the same technique can shrink 3D CAD programs (Section 6.1), make floating point more accurate (Section 6.2), optimize deep learning compute graphs (Section 6.3), and synthesize the very rules that it needs to work (Section 6.4), then why wouldn’t it work for your problem in your domain? Even outside of those case studies, *egg* is powering or inspiring many additional projects, only some of which are published at the time of writing this [WHL⁺20, Che21, VNL⁺21].

People used and worked on equality saturation before I began work on this thesis. In fact, I first learned about it by talking to some of those people (the Herbie developers, Section 6.2). My first reaction to equality saturation was that of a skeptic (“it’s just unionfind”), and I began work on *egg* out of hubris to show that it was trivial to implement. Predictably, I was quickly humbled, but eventually *egg* became a useful tool. After some more learning, I moved into the fanatic phase: “why isn’t everyone doing this?” Through the course of pushing a round *egg* through several square holes, I would like to think that I have adopted a more pragmatic approach,¹ informed by the strengths and weaknesses of the technique.

¹Friends and colleagues, however, will be excused for still thinking me a fanatic. 🤖

This work in this thesis puts a dent in those weaknesses and introduces some new strengths. With these advances, and with egg packaging them all up, I hope that others find that equality saturation is a viable, useful, and even fun technique with a more straightforward journey than my own.

7.1 Future Work

Going forward, I hope and expect that equality saturation will take a larger role in all kinds of compilers, synthesizers, and optimizers. But further work is needed to get there. While egg can already be used to build a state-of-the-art optimizer fast enough to use in a compiler [YPW⁺21], this only applies for domains with algebraic, context-free interpretations of their expressions. To make equality saturation practical in more complex domains (like general purpose programming languages), new conceptual advances are required.

Prior work [TSTL09] introduced *Program Expression Graphs* (PEGs) to represent programs with mutation and loops inside e-graphs, but not source-level programs or functional IRs. These latter two are tricky for e-graphs because binding means that equivalence is *contextual*: two variables x in the program might refer to different binding sites. E-class analyses could be combined with recent work [MEL⁺21] that proposes a modular way to describe binding structure.

Complex domains will require larger search spaces, further pushing the performance requirements for e-graphs and equality saturation. This work introduced a new, more efficient algorithm for congruence closure in e-graphs; the remaining bottleneck is e-matching. egg uses the current state-of-the-art [dMB07] backtracking algorithm that wastes work when search for patterns like $(x * y) + (x * z)$ with multiple occurrences of the same variables. A smarter e-matching algorithm could take advantage of these equality constraints, similar to how joins in a relational database do.

Finally, equality saturation’s performance and correctness both hinge on the rewrites provided

by the users. If those rewrites are incomplete or unsound, equality saturation may miss optimizations or yield incorrect results. Automatically generating these rules could make using equality saturation even easier for prospective users.

Bibliography

- [ABC⁺16] Martín Abadi, Paul Barham, Jianmin Chen, Zhifeng Chen, Andy Davis, Jeffrey Dean, Matthieu Devin, Sanjay Ghemawat, Geoffrey Irving, Michael Isard, Manjunath Kudlur, Josh Levenberg, Rajat Monga, Sherry Moore, Derek G. Murray, Benoit Steiner, Paul Tucker, Vijay Vasudevan, Pete Warden, Martin Wicke, Yuan Yu, and Xiaoqiang Zheng. TensorFlow: A System for Large-Scale Machine Learning. In *12th USENIX Symposium on Operating Systems Design and Implementation (OSDI 16)*, 2016.
- [AEH⁺99] Marc Andries, Gregor Engels, Annegret Habel, Berthold Hoffmann, Hans-Jörg Krewowski, Sabine Kuske, Detlef Plump, Andy Schürr, and Gabriele Taentzer. Graph transformation for specification and programming. *Sci. Comput. Program.*, 34(1):1–54, April 1999.
- [BA06] Sorav Bansal and Alex Aiken. Automatic generation of peephole superoptimizers. In *Proceedings of the 12th International Conference on Architectural Support for Programming Languages and Operating Systems, ASPLOS XII*, page 394–403, New York, NY, USA, 2006. Association for Computing Machinery.
- [BC10] Yves Bertot and Pierre Castran. *Interactive Theorem Proving and Program Development: Coq’Art The Calculus of Inductive Constructions*. Springer Publishing Company, Incorporated, 1st edition, 2010.
- [BCD⁺11] Clark Barrett, Christopher L. Conway, Morgan Deters, Liana Hadarean, Dejan Jovanović, Tim King, Andrew Reynolds, and Cesare Tinelli. Cvc4. In *Proceedings of the 23rd International Conference on Computer Aided Verification, CAV’11*, page 171–177, Berlin, Heidelberg, 2011. Springer-Verlag.
- [Bli13] Gabriel Hjort Blindell. Survey on instruction selection: An extensive and modern literature review. *CoRR*, abs/1306.4898, 2013.
- [BRTV00] Leo Bachmair, IV Ramakrishnan, Ashish Tiwari, and Laurent Vigneron. Congruence closure modulo associativity and commutativity. In *International Workshop on Frontiers of Combining Systems*, pages 245–259. Springer, 2000.

- [Che21] Alessandro Cheli. Metatheory.jl: Fast and elegant algebraic computation in julia with extensible equality saturation. *Journal of Open Source Software*, 6(59):3078, 2021.
- [CMJ⁺18] Tianqi Chen, Thierry Moreau, Ziheng Jiang, Lianmin Zheng, Eddie Yan, Haichen Shen, Meghan Cowan, Leyuan Wang, Yuwei Hu, Luis Ceze, Carlos Guestrin, and Arvind Krishnamurthy. TVM: An automated end-to-end optimizing compiler for deep learning. In *13th USENIX Symposium on Operating Systems Design and Implementation (OSDI 18)*, pages 578–594, Carlsbad, CA, October 2018. USENIX Association.
- [DCLT19] J. Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding. In *NAACL-HLT*, 2019.
- [Der93] Nachum Dershowitz. A taste of rewrite systems. In Peter E. Lauer, editor, *Functional Programming, Concurrency, Simulation and Automated Reasoning: International Lecture Series 1991–1992 McMaster University, Hamilton, Ontario, Canada*, pages 199–228. Springer Berlin Heidelberg, Berlin, Heidelberg, 1993.
- [DJ90] Nachum Dershowitz and Jean-Pierre Jouannaud. Rewrite systems. In Jan van Leeuwen, editor, *Handbook of Theoretical Computer Science, Volume B: Formal Models and Semantics*, pages 243–320. Elsevier and MIT Press, 1990.
- [dMB07] Leonardo de Moura and Nikolaj Bjørner. Efficient e-matching for smt solvers. In Frank Pfenning, editor, *Automated Deduction – CADE-21*, pages 183–198, Berlin, Heidelberg, 2007. Springer Berlin Heidelberg.
- [DMB08] Leonardo De Moura and Nikolaj Bjørner. Z3: An efficient smt solver. In *Proceedings of the Theory and Practice of Software, 14th International Conference on Tools and Algorithms for the Construction and Analysis of Systems, TACAS’08/ETAPS’08*, pages 337–340, Berlin, Heidelberg, 2008. Springer-Verlag.
- [DNS05] David Detlefs, Greg Nelson, and James B. Saxe. Simplify: A theorem prover for program checking. *J. ACM*, 52(3):365–473, May 2005.
- [DP60] Martin Davis and Hilary Putnam. A computing procedure for quantification theory. *J. ACM*, 7(3):201–215, July 1960.
- [DPIP⁺18] Tao Du, Jeevana Priya Inala, Yewen Pu, Andrew Spielberg, Adriana Schulz, Daniela Rus, Armando Solar-Lezama, and Wojciech Matusik. Inversecsg: automatic conversion of 3d models to csg trees. In *ACM Transactions on Graphics*, pages 1–16, 12 2018.
- [DST80] Peter J. Downey, Ravi Sethi, and Robert Endre Tarjan. Variations on the common subexpression problem. *J. ACM*, 27(4):758–771, October 1980.

- [ERSLT18] Kevin Ellis, Daniel Ritchie, Armando Solar-Lezama, and Joshua B. Tenenbaum. Learning to infer graphics programs from hand-drawn images. In *Neural Information Processing Systems (NIPS)*, 2018.
- [FSWC20] Jingzhi Fang, Yanyan Shen, Yue Wang, and Lei Chen. Optimizing DNN Computation Graph Using Graph Substitutions. *Proc. VLDB Endow.*, 13(12):2734–2746, July 2020.
- [GAB⁺20] Gerald Gamrath, Daniel Anderson, Ksenia Bestuzheva, Wei-Kun Chen, Leon Eifler, Maxime Gasse, Patrick Gemander, Ambros Gleixner, Leona Gottwald, Katrin Halbig, Gregor Hendel, Christopher Hojny, Thorsten Koch, Pierre Le Bodic, Stephen J. Maher, Frederic Matter, Matthias Miltenberger, Erik Mühmer, Benjamin Müller, Marc E. Pfetsch, Franziska Schlösser, Felipe Serrano, Yuji Shinano, Christine Tawfik, Stefan Vigerske, Fabian Wegscheider, Dieter Weninger, and Jakob Witzig. The SCIP Optimization Suite 7.0. Technical report, Optimization Online, March 2020.
- [Goo17] Google. XLA - TensorFlow, compiled, Mar 2017.
- [HA00] James C. Hoe and Arvind. *Hardware Synthesis from Term Rewriting Systems*, pages 595–619. Springer US, Boston, MA, 2000.
- [Her21] Herbie. Herbie can generate more-complex expressions that aren’t more precise, 2021. <https://github.com/uwplse/herbie/issues/261#issuecomment-680896733>.
- [IMA⁺17] Forrest N. Iandola, Matthew W. Moskewicz, K. Ashraf, Song Han, W. Dally, and K. Keutzer. Squeezenet: Alexnet-level accuracy with 50x fewer parameters and <1mb model size. *ArXiv*, abs/1602.07360, 2017.
- [JNR02] Rajeev Joshi, Greg Nelson, and Keith Randall. Denali: A goal-directed superoptimizer. *SIGPLAN Not.*, 37(5):304–314, May 2002.
- [JPT⁺19] Zhihao Jia, Oded Padon, James Thomas, Todd Warszawski, Matei Zaharia, and Alex Aiken. TASO: Optimizing Deep Learning Computation with Automatic Generation of Graph Substitutions. In *Proceedings of the 27th ACM Symposium on Operating Systems Principles, SOSP ’19*, page 47–62, New York, NY, USA, 2019. Association for Computing Machinery.
- [JTW⁺19] Zhihao Jia, James Thomas, Todd Warszawski, Mingyu Gao, Matei Zaharia, and Alex Aiken. Optimizing DNN Computation with Relaxed Graph Substitutions. In *Proceedings of the 2nd SysML Conference, SysML ’19*, 2019.
- [LAB⁺20] Chris Lattner, Mehdi Amini, Uday Bondhugula, Albert Cohen, Andy Davis, Jacques Pienaar, River Riddle, Tatiana Shpeisman, Nicolas Vasilache, and Oleksandr Zinenko. MLIR: A compiler infrastructure for the end of moore’s law, 2020.

- [LD15] S. Liu and W. Deng. Very deep convolutional neural network based image classification using small training sample size. In *2015 3rd IAPR Asian Conference on Pattern Recognition (ACPR)*, pages 730–734, 2015.
- [Mas87] Henry Massalin. Superoptimizer: A look at the smallest program. In *Proceedings of the Second International Conference on Architectural Support for Programming Languages and Operating Systems, ASPLOS II*, page 122–126, Washington, DC, USA, 1987. IEEE Computer Society Press.
- [MEL⁺21] Krzysztof Maziarczyk, Tom Ellis, Alan Lawrence, Andrew Fitzgibbon, and Simon Peyton Jones. Hashing modulo alpha-equivalence. In *ACM SIGPLAN Conference on Programming Language Design and Implementation (PLDI’21)*. ACM, ACM, June 2021.
- [MN17] David Menendez and Santosh Nagarakatte. Alive-infer: Data-driven precondition inference for peephole optimizations in llvm. In *Proceedings of the 38th ACM SIGPLAN Conference on Programming Language Design and Implementation, PLDI 2017*, page 49–63, New York, NY, USA, 2017. Association for Computing Machinery.
- [Nel80] Charles Gregory Nelson. *Techniques for Program Verification*. PhD thesis, Stanford University, Stanford, CA, USA, 1980. AAI8011683.
- [NJK⁺20] Julie L. Newcomb, Steven Johnson, Shoaib Kamil, Andrew Adams, and Ratislav Bodik. Verifying and improving halide’s term rewriting system with program synthesis. *Proceedings of the ACM on Programming Languages (OOPSLA)*, 2020.
- [NO05] Robert Nieuwenhuis and Albert Oliveras. Proof-producing congruence closure. In *Proceedings of the 16th International Conference on Term Rewriting and Applications, RTA’05*, page 453–468, Berlin, Heidelberg, 2005. Springer-Verlag.
- [NRB⁺19] Andres Nötzli, Andrew Reynolds, Haniel Barbosa, Aina Niemetz, Mathias Preiner, Clark Barrett, and Cesare Tinelli. Syntax-guided rewrite rule enumeration for smt solvers. In Mikoláš Janota and Inês Lynce, editors, *Theory and Applications of Satisfiability Testing – SAT 2019*, pages 279–297, Cham, 2019. Springer International Publishing.
- [NRR14] Hung Q Ngo, Christopher Ré, and Atri Rudra. Skew strikes back: New developments in the theory of join algorithms. *SIGMOD Rec.*, 42(4):5–16, February 2014.
- [NVI] NVIDIA. TensorRT.
- [NWA⁺20] Chandrakana Nandi, Max Willsey, Adam Anderson, James R. Wilcox, Eva Darulova, Dan Grossman, and Zachary Tatlock. Synthesizing structured CAD models with equality saturation and inverse transformations. In *Proceedings of the 41st ACM*

- SIGPLAN Conference on Programming Language Design and Implementation*, PLDI 2020, page 31–44, New York, NY, USA, 2020. Association for Computing Machinery.
- [NWP02] Tobias Nipkow, Markus Wenzel, and Lawrence C. Paulson. *Isabelle/HOL: A Proof Assistant for Higher-Order Logic*. Springer-Verlag, Berlin, Heidelberg, 2002.
- [NWP⁺18] Chandrakana Nandi, James R. Wilcox, Pavel Panchekha, Taylor Blau, Dan Grossman, and Zachary Tatlock. Functional programming for compiling and decompiling computer-aided design. *Proc. ACM Program. Lang.*, 2(ICFP):99:1–99:31, July 2018.
- [NWZ⁺21] Chandrakana Nandi, Max Willsey, Amy Zhu, Brett Saiki, Yisu Wang, Adam Anderson, Adriana Schulz, Dan Grossman, and Zachary Tatlock. Rewrite rule inference using equality saturation, 2021. Under submission.
- [PF] Laurent Perron and Vincent Furnon. Or-tools.
- [PGM⁺19] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Kopf, Edward Yang, Zachary DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. PyTorch: An imperative style, high-performance deep learning library. In H. Wallach, H. Larochelle, A. Beygelzimer, F. d'Alché-Buc, E. Fox, and R. Garnett, editors, *Advances in Neural Information Processing Systems 32*, pages 8024–8035. Curran Associates, Inc., 2019.
- [PJTH01] Simon Peyton Jones, Andrew Tolmach, and Tony Hoare. Playing by the rules: rewriting as a practical optimisation technique in ghc. In *2001 Haskell Workshop*. ACM SIGPLAN, September 2001.
- [PKSL20] Varot Premtoon, James Koppel, and Armando Solar-Lezama. Semantic code search via equational reasoning. In *Proceedings of the 41st ACM SIGPLAN Conference on Programming Language Design and Implementation*, PLDI 2020, page 1066–1082, New York, NY, USA, 2020. Association for Computing Machinery.
- [PSSWT15] Pavel Panchekha, Alex Sanchez-Stern, James R. Wilcox, and Zachary Tatlock. Automatically improving accuracy for floating point expressions. *SIGPLAN Not.*, 50(6):1–11, June 2015.
- [RKBA⁺13] Jonathan Ragan-Kelley, Connelly Barnes, Andrew Adams, Sylvain Paris, Frédo Durand, and Saman Amarasinghe. Halide: A language and compiler for optimizing parallelism, locality, and recomputation in image processing pipelines. In *Proceedings of the 34th ACM SIGPLAN Conference on Programming Language Design and Implementation*, PLDI '13, page 519–530, New York, NY, USA, 2013. Association for Computing Machinery.

- [Rus] Rust. Rust programming language. <https://www.rust-lang.org/>.
- [SGL⁺17] Gopal Sharma, Rishabh Goyal, Difan Liu, Evangelos Kalogerakis, and Subhransu Maji. Csgnet: Neural shape parser for constructive solid geometry. *CoRR*, abs/1712.08290, 2017.
- [SSL16] Rohit Singh and Armando Solar-Lezama. Swapper: A framework for automatic generation of formula simplifiers based on conditional rewrite rules. In *Proceedings of the 16th Conference on Formal Methods in Computer-Aided Design*, FMCAD ’16, page 185–192, Austin, Texas, 2016. FMCAD Inc.
- [STL11] Michael Stepp, Ross Tate, and Sorin Lerner. Equality-based translation validator for llvm. In Ganesh Gopalakrishnan and Shaz Qadeer, editors, *Computer Aided Verification*, pages 737–742, Berlin, Heidelberg, 2011. Springer Berlin Heidelberg.
- [SVI⁺16] Christian Szegedy, V. Vanhoucke, S. Ioffe, Jon Shlens, and Z. Wojna. Rethinking the inception architecture for computer vision. *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 2818–2826, 2016.
- [Tar75] Robert Endre Tarjan. Efficiency of a good but not linear set union algorithm. *J. ACM*, 22(2):215–225, April 1975.
- [TLS⁺19] Yonglong Tian, Andrew Luo, Xingyuan Sun, Kevin Ellis, William T. Freeman, Joshua B. Tenenbaum, and Jiajun Wu. Learning to infer and execute 3d shape programs. In *International Conference on Learning Representations*, 2019.
- [TSTL09] Ross Tate, Michael Stepp, Zachary Tatlock, and Sorin Lerner. Equality saturation: A new approach to optimization. In *Proceedings of the 36th Annual ACM SIGPLAN-SIGACT Symposium on Principles of Programming Languages*, POPL ’09, pages 264–276, New York, NY, USA, 2009. ACM.
- [VNL⁺21] Alexa VanHattum, Rachit Nigam, Vincent T. Lee, James Bornholt, and Adrian Sampson. Vectorization for digital signal processors via equality saturation. In *Proceedings of the 26th ACM International Conference on Architectural Support for Programming Languages and Operating Systems*, ASPLOS 2021, page 874–886, New York, NY, USA, 2021. Association for Computing Machinery.
- [WHL⁺20] Yisu Remy Wang, Shana Hutchison, Jonathan Leang, Bill Howe, and Dan Suciu. SPORES: Sum-product optimization via relational equality saturation for large scale linear algebra. *Proceedings of the VLDB Endowment*, 2020.
- [WS97] Deborah L. Whitfield and Mary Lou Soffa. An approach for exploring code improving transformations. *ACM Trans. Program. Lang. Syst.*, 19(6):1053–1084, November 1997.

- [WWF⁺20] Max Willsey, Yisu Remy Wang, Oliver Flatt, Chandrakana Nandi, Pavel Panchekha, and Zachary Tatlock. egg: Fast and extensible e-graphs, 2020.
- [WZN⁺19] Chenming Wu, Haisen Zhao, Chandrakana Nandi, Jeffrey I. Lipton, Zachary Tatlock, and Adriana Schulz. Carpentry compiler. *ACM Transactions on Graphics*, 38(6):Article No. 195, 2019. presented at SIGGRAPH Asia 2019.
- [XGD⁺17] S. Xie, R. Girshick, P. Dollár, Z. Tu, and K. He. Aggregated residual transformations for deep neural networks. In *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 5987–5995, 2017.
- [YPW⁺21] Yichen Yang, Phitchaya Mangpo Phothilimtha, Yisu Remy Wang, Max Willsey, Sudip Roy, and Jacques Pienaar. Equality saturation for tensor graph superoptimization. In *Proceedings of Machine Learning and Systems*, 2021.
- [ZL17] Barret Zoph and Quoc V. Le. Neural architecture search with reinforcement learning. *arXiv*, 2017.
- [ZVSL18] B. Zoph, V. Vasudevan, J. Shlens, and Q. V. Le. Learning transferable architectures for scalable image recognition. In *2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 8697–8710, 2018.