

Automatic Incrementalization of Vertex-Centric Programs

Timothy A. K. Zakian
University of Oxford
timothy.zakian@cs.ox.ac.uk

Ludovic Capelli
University of Edinburgh
L.Capelli@ed.ac.uk

Zhenjiang Hu
National Institute of Informatics
hu@nii.ac.jp

ABSTRACT

Vertex-centric computations by their very definition are communication intensive. Because of this, a large portion of the total computational time of a vertex-centric program is spent in communication, and reducing the number of messages sent in order to perform a given computation is a critical aspect of any optimization efforts. While previous work has focused on reducing communication overhead by directly changing the communication patterns by altering the way the graph is partitioned and distributed, or by altering the graph topology itself, in this paper we present a different optimization strategy based on a family of complementary *compile-time program transformations* to minimize communication overhead. In particular, we show how through a series of compile-time transformations of the program we can automatically *incrementalize* the vertex-centric algorithm as a whole. We formalize and empirically evaluate these transformations and show that they can result in up-to 4.4X speedup on certain programs. Furthermore, due to the fact that these are program transformations alone, other prior optimization strategies can work with the resulting vertex-centric program just as they would a non-incrementalized program.

CCS CONCEPTS

• **Software and its engineering** → **Distributed programming languages**; **Compilers**; **Domain specific languages**; • **Hardware** → Emerging languages and compilers;

KEYWORDS

Pregel, big graph processing, domain specific languages, incrementalization.

ACM Reference Format:

Timothy A. K. Zakian, Ludovic Capelli, and Zhenjiang Hu. 2018. Automatic Incrementalization of Vertex-Centric Programs. In *Proceedings of ACM International Conference on Parallel Programming (ICPP '18)*. ACM, New York, NY, USA, Article 4, 10 pages. https://doi.org/10.475/123_4

1 INTRODUCTION

As the scale of real-world graphs that need to be computed over rapidly grows, with graphs many times being made of billions of nodes and edges, the need for *highly-parallel*, and *scalable* graph processing solutions has grown ever greater. One such model for computing over large-scale graphs is the *vertex-centric* model [13]. The vertex-centric model of graph computations has been the center of an intense amount of recent research and development efforts, and has lead to popular vertex-centric frameworks such as Pregel [12], Pregel+ [19], GraphLab [11], Apache Giraph [6] and many more [2,

9, 16]. These vertex-centric programming systems have allowed programmers to easily express computations over graphs at the largest scales [3, 12], and have become not just a useful, but crucial part of a programmer’s arsenal when working with graphical data at scale.

Each of these modern vertex-centric frameworks are based on the Bulk Synchronous Parallel (BSP) model of computation [18], and the computation is performed as a series of iterative *supersteps*. Each superstep is comprised of a parallel execution of a user-defined program—which takes in messages from other vertices along incoming edges and sends messages along outgoing edges—over the vertices of the graph as they are distributed among computation nodes or *workers*, followed by communication between vertices, and finally a global barrier synchronization is performed to ensure that all messages have been delivered before the start of the next superstep.

The vertex-centric model of computation sacrifices memory and computational locality for scalability, thus while the model is incredibly scalable it is by its very nature incredibly communication intensive. Due to this inherent communication overhead, since the original Pregel paper [12] popularized the computational model there has been an explosion of research in trying to speedup computations in this setting by hybridizing the model to be a subgraph-centric computational model [7, 17], and determine better partitioning of the graph amongst the computational nodes [4, 19]. All with the eventual goal of minimizing the number of messages sent in the computation as a whole, and minimizing communication bottlenecks. This is with good reason: the vertex-centric model is so communication heavy that many times the majority of time taken by the computations is spent performing communication.

The techniques that we present in this paper to incrementalize and ensure “meaningful-only” messages can be seen as being motivated by this same view towards optimizing vertex-centric programs by minimizing communication. However, while previous research has sought to minimize communication by changing the computational model visible to the programmer, or by changing the graphical structure of the program, we take a different view. Particularly, we guide our research quest by the following question:

How do we make sure that every message that is sent during the computation is *meaningful*?

Where we intuitively call a message sent from a vertex meaningful if it differs from the most recently sent message from that vertex.

Once our program has this property where every message sent by a vertex is meaningful, we can *incrementalize* the computation as a whole: since messages are only sent when a message value has changed, each “whole message” can be converted to a “ Δ -message”, where instead of sending the value for a message, the delta from the previous message is sent instead. In this setting a lack of a message from a vertex u represents an affirmation of cache coherency with u , and a Δ -message from u represents a “patch” to be applied to the cache on the receiving vertex in order to keep it coherent with the updated message value from u .

However, incrementalizing the computation in this manner requires that we change the way the computation within each vertex interacts with the messages it receives. We therefore need to specify a number of transformations that need to be performed in order to allow not just sending meaningful-only messages, but also sending and interacting with Δ -messages. Because of this, we need to depart from the majority of previous work which has taken a library-oriented or *shallow embedding* view towards vertex-centric computing and instead move to a deeply-embedded setting where we have access to the program AST and can easily analyze and manipulate it.

Changing the world in which vertex-centric computations are expressed from the shallowly-embedded computational-model-as-library realm into one in which users write programs that are *compiled* down to a vertex-centric framework such as Pregel+ means that our guiding question has now changed somewhat:

What *compile-time transformations* can be performed in order to ensure that every message sent by the compiled program is meaningful?

Since we are compiling to a vertex-centric framework—Pregel+—the optimizations and transformations that we perform in order to ensure the above property could be obtained by-hand by anyone in any of the library-based frameworks we have mentioned thus far. However, taking the language-based approach presented in this paper, the programmer does not have to think about any of these optimizations themselves when writing their program and instead writes a normal vertex-centric algorithm, which is both shorter, easier to understand, boiler-plate free, and faster than the equivalent non-“ninja”d algorithm in Pregel+.

We make the following contributions.

- We present a novel message passing policy for vertex-centric computations that ensures that only new (“meaningful”) results are communicated, and where the default state for a vertex is to be halted.
- We present the first vertex-centric language with automatic incrementalization of messages and vertex-computations.
- We evaluate the suite of transformations that we present, and show that they can reduce execution time by up-to 4.4X and messages up-to 5.8X.
- We evaluate programs in which incrementalization is *not* useful, and show that our transformations have not slowed down compiled programs; thus these transformations can always be performed without penalty.

2 VERTEX-CENTRIC COMPUTATION

As mentioned in Section 1 Pregel-like vertex-centric computational frameworks are designed around the BSP model of computation, where different vertices of the graph are distributed or associated to different machines in the cluster, and where each vertex u keeps track of its adjacency list (*i.e.*, the neighbors of u) along with a *vertex state* which can contain any number of user-defined fields that the user can access, and whose values persist from one superstep to another. In order to define a computation over a graph in this model, the programmer implements a `compute()` function, which runs on each vertex in the graph and proceeds in iterations called *supersteps*. In each superstep the `compute()` function is called on each currently *active* vertex

u in the graph.¹ The `compute()` function then performs the task specified by the user for each active vertex u . Each `compute()` function has access to the messages from incoming edges at the previous superstep, can read and write to its own vertex state, can send messages which are to be received in the next superstep to the vertex’s neighbors, and can make the vertex inactive by calling a `vote_to_halt()` function. Once a vertex is halted, it is no longer considered for computation unless it is “reactivated” by receiving an incoming message. Once all vertices in a computation have halted, and there are no more messages pending, the computation as a whole terminates.

Since the computation is distributed amongst many different machines m_i with often-times multiple workers per-machine, Pregel(+) allows users to implement message *combiners* which combine messages from (all the vertices on) one machine m_j to another vertex u on a different machine m_i . These combiners are crucial to reducing inter-worker communication; by using them instead of sending n messages from a one machine to a single vertex, we can instead combine all of the n local messages on the machine into a single message before sending that *single* message to the destination vertex. However, since there is no guarantee about the order in which messages are combined, the combiner operation is restricted to be both commutative and associative.

For a more detailed and comprehensive description of the programming model and the various aspects and features of it, we refer the reader to Malewicz et al. [12].

3 BUILDING SOME INTUITION

Many times during a vertex-centric program, a given vertex may compute the same value to send as a message from one superstep to the next. We’ll call such messages “meaningless”: a repeated message with the same value does not provide any new information to the computation that could not be remembered from previous (vertex-local) information². This section provides the intuition and motivation for the transformations that we will be performing in order to prevent sending meaningless messages and to incrementalize the compiled program. We’ll use PageRank [14] as it is defined in Figure 1 as a running example throughout this section.

3.1 Meaningful Messages

We can imagine many cases in which the `compute()` function calculates the same pagerank value (`msg`) from one superstep to another. In order to make sure that we don’t send the same message twice in a row, we modify the code to check if the value of `msg` has changed from the previous superstep before sending that value. We thus replace lines 15-17 of Figure 1 with the following piece of code:

```
...
// Determine if we have a new value to send
if (msg != value().old_msg) {
    for (VertexID v : value().neighbors) {
        send_message(v, msg); // New value, so send
    }
    // Update most recently sent value
    value().old_msg = msg;
} // Otherwise, don't send
```

¹At the beginning of the computation all vertices start out *active*.

²In other words, by accurately *memoizing* the vertex computation.

```

1 void compute(const MessageContainer& messages) {
2     if(step_num() == 1) {
3         value().pr = 1.0 / get_vnum();
4     } else {
5         double sum = 0;
6         for(double message : messages) {
7             sum += message;
8         }
9         value().pr = 0.15 + 0.85 *
10            (sum / get_vnum());
11     }
12     if(step_num() < 30) {
13         double msg = value().pr /
14            value().neighbors.size();
15         for (VertexID v : value().neighbors) {
16             send_message(v, msg);
17         }
18     } else { vote_to_halt(); }
19 }

```

Figure 1: The compute() function for PageRank in Pregel+.

...

Where we remember the old value of msg in a new vertex field value().old_msg. However after performing this transformation we lose a key invariant of the computation: previously each vertex received *all* of its neighbors pageranks at *every* superstep, however in the transformed program, messages are no longer sent unless the message value has changed. Because of this, the summation of the neighbors pageranks in lines 6-8 of Figure 1 is no longer correct since (possibly) not all of the neighbors pageranks will be sent at every superstep. We therefore need change how we interact with the messages sent by our neighbors, and change the way we calculate their sum.

3.2 Memoizing Aggregations

In order to get around this problem, we memoize the summation of our neighbors pageranks from one superstep to the next. This is possible since while we have lost the invariant that we receive the pagerank for all of our neighbors, we have gained the invariants that

- (1) If we receive a message it is non-equal to the previous message sent from that vertex; and
- (2) If we haven't recieved a message from a vertex, then the value of the message it *would* have sent is the same as the most recent message from that vertex.

Using this knowledge, we can memoize the summation from one superstep to the next by adding a field to our vertex—value().sum—that records the sum computed at the previous superstep. This field is then used (and updated) to get the sum for the current superstep once we recieve a set of messages:

...

```

for (double message : messages) {
    value().sum += message;
}
...

```

But there's a problem with this method: we still don't calculate the correct sum. While we are only sending changing values, we are sending *whole messages* i.e., our messages carry the *new* pagerank for a vertex as opposed to the *delta from the previous pagerank*. Thus when we memoize the aggregation as above, we wind up double-counting the previous values for the message that we received in the sum. In order to fix this, we need to change the values that we send between vertices, from *whole messages* to Δ -messages; instead of sending a value from which we can compute the new value, we send a value that allows us to compute the *change* to the already computed value in order to arrive at the correct new value.

3.3 Δ -Messages

As we have just seen, in order to only send meaningful messages, we need to also memoize the aggregation operations operating over the messages coming in. However in order to allow this inner memoization of our vertex computations, we need to change the messages that we send: from each neighbor sending its new value, to each neighbor sending its "delta" from its previous value. e.g., for the running example, if on superstep n a vertex computed msg = 0.001 and on superstep $n + 1$ it computed msg = 0.02 it would send a Δ -message of +0.019 instead of 0.02.

To send modified messages, we introduce a function computeDelta that given an old message value, and a new message value computes the delta between the two:

```

double computeDelta(double oldMsg, double newMsg) {
    return newMsg - oldMsg;
}

```

Calls to send_message are then updated so instead of directly sending msg as was done previously, the delta between the old and new message using computeDelta is computed and then the resulting Δ -message is sent. Thus the final code now for lines 15-17 in Figure 1 is

```

...
for (VertexID v : value().neighbors) {
    // Determine if we have a new value to send
    if (msg != value().old_msg) {
        // New value, so compute the delta
        double delta = computeDelta(value().old_msg, msg);
        send_message(v, delta); // send the _delta_
        // Update most recently sent value
        value().old_msg = msg;
    } // Otherwise, don't send
}
...

```

At this point, our transformed program has the following new properties:

- (1) We only send a message when the value being sent is a meaningful update;
- (2) The "boundary" of the vertex-function has been incrementalizes so it only aggregates over the set of changed values at each step instead of all neighbors; and
- (3) The messages that we send represent *changes* to previous message values as opposed to *new message values*.

It is important to realize the dependencies between each of the program transformations that have just been performed in order to

arrive at these properties; memoization (incrementalization) could not be performed without incrementalizing the messages that are sent, and both of these optimizations are only useful in the presence of a meaningful-only messaging policy.

The series of program transformations that we have performed to our PageRank program, can be mechanized as compiler passes. Furthermore, the memoization and incrementalization transformations can be generalized to arbitrary commutative and associative operations. In the rest of this paper we formalize the properties about vertex-centric computations, messages, and program transformations that we have discussed in this section and demonstrate the usefulness of the program transformations that we have just introduced empirically.

4 MAKING THINGS PRECISE

Now that we have built our intuition about why, and how we transform programs we turn our attention to formalizing these notions.

4.1 The Meaning of Meaningful Messages

Earlier on we appealed to the reader's intuition regarding what constitutes a "meaningful" message. We now make this notion precise.

Definition 1 (Meaningful Messages). *A message m_1 sent from a vertex u_1 to another vertex u_2 at superstep n , is meaningful if any of the following statements hold:*

- (1) $n = 0$; or
- (2) If m_2 is the most recent message sent from u_1 to u_2 , then $m_1 \neq m_2$; or
- (3) No message was sent from u_1 to u_2 for all supersteps $k, k < n$.

In other words, at the beginning of the computation every message is considered "meaningful", and after that we never want to send two messages with the same value in a row between the same two vertices. This definition of what constitutes a meaningful message is meant to encapsulate the idea of memoization in the vertex-centric context: a message is meaningful iff there is no efficient way for the other (receiving) vertex to cache that messages value from the most recent value sent by that vertex.

In this meaningful-messages only setting, the absence of a message to a vertex is meaningful in itself, since it tells us that the previous value of the *sending* vertex has not changed i.e., that the *cache* is *coherent* for that vertex (as viewed from the receiving vertex), and we therefore do not need to invalidate the local "cache" for that vertex's value on the receivers side. Thus in this setting messages between vertices not only represent new values to be used, but also in a sense *cache invalidations* for the receiving vertex's cache. However, due to the distributed nature of the computation, how we update that cache after we have "invalidated" it with a new message is a crucial. This is where an *incrementalization policy* is brought into play.

4.2 Incrementalization

4.2.1 Memoization: The Inefficient Approach. When we think of caching for vertices, the first idea that comes to mind is to cache the individual value for each of the vertex's neighbors locally within the vertex's state based on the approach in Figure 2b where we keep a lookup table from vertex id to that vertex's value. When we receive a message we update this lookup table with the new mapping from the sending vertex's identifier to value, and instead of directly

iterating over the messages within the vertex function, we instead use this lookup table as a proxy for the messages to that vertex.

While this lookup-table-based method of caching neighboring vertices message values allows us to fulfil the requirement of only sending meaningful messages throughout the computations, it has a number of issues that make this method impractical and inefficient. In particular: each message sent must be tagged with the sending vertex's id which in many cases can double the size of each message; and each vertex keeping a local lookup table from vertex id to value can increase the size of the vertex state and hence memory footprint of the computation considerably.

As the reader has probably surmised after reading Section 3, this is not the approach we take in this paper. While we are able to get rid of meaningless messages through the above memoization this reduction in messaging comes at such a considerable cost that the resulting computation can run even slower than the original. The key realization to resolving these issues lies in blurring the interaction between the messages we receive and the computation that we perform within the vertex function. In other words, we don't want to simply cache neighboring message values, but instead *incrementalize* the *interactions* between vertices.

4.2.2 Efficiency Through Incrementalization. Automatically incrementalizing programs is a well-studied and active field of research, however in this setting we are only interested in methods which allow us to *statically* incrementalize the program using compile-time program transformations. While a standard way of performing static incrementalization of programs is by deriving a transformation ∂ [1] such that given a value a , a function f and a change from that a , da , the following equation holds³

$$f(a \oplus da) \approx (fa) \oplus (\partial f a da) \quad (1)$$

However, in our case this version of incrementalization—while static and only using program transformations—will not work since our goal is to not necessarily incrementalize the vertex-program itself, but instead incrementalize the vertex's *interactions with other vertices* and through this incrementalize the vertex computation as a whole.

Beyond simply needing to incrementalize the interactions as opposed to the vertex-program itself, we also want this transformation to be subject to the conditions that the resulting transformed program is still efficient: the resulting program should use a minimal amount of additional memory per vertex compared to the non-incrementalized program, and the size of the messages should not change.⁴ Furthermore, due to the highly distributed nature of vertex-centric computations, given a new message m' to send, and the previous message m we need to be able to determine the incrementalized, or Δ -message, for m' on the senders side (i.e., only using the senders vertex state).⁵ This locally determinable incrementalization of messages is a crucial property since it helps ensure the efficiency of the incrementalization policy; without this property we would incur massive communication overheads and have to impose inefficient caching policies similar to the one discussed in Section 4.2.1. Thus as we will see in Sections 6.4 and 6.5 the compile-time transformations performed to the computation *within*

³Where $\approx$ is meant to denote denotational equality of expressions, and $\oplus$ is a binary operation subject to certain conditions [1].

⁴Or change as little as possible. We discuss this more in Section 9.

⁵If a message is not a Δ -message, we will call it a "full message".

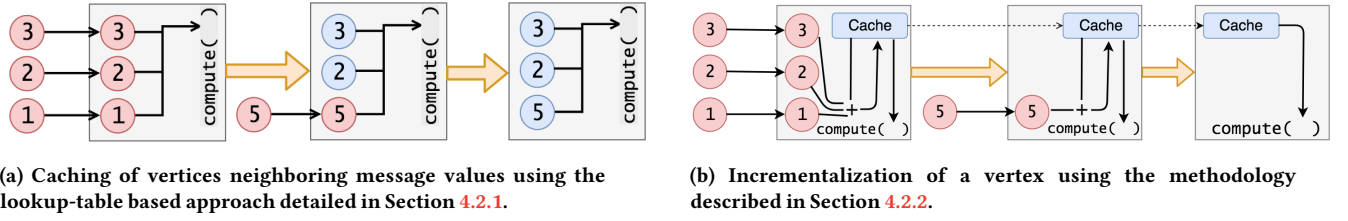


Figure 2: Different (vertex-internal) caching strategies for vertex-centric computation. Blue vertices are cached, while red vertices have just been updated.

```

u      ::= vertexName
uop    ::= - | not
pop    ::= min | max
op     ::= + | - | * | / | &&
        | || | < | > | ≥ | ≤ | ==
⊞     ::= + | * | min | max | || | &&
g      ::= #in | #out | #neighbors
a      ::= ref | xf
τ      ::= int | bool | float
prog   ::= init {e}; stmt
stmt   ::= step {e} | iter i {e} until {e} | stmt; stmt
e      ::= x | xf | b | i | f | e op e | pop e1 e2 | uop e
        | graphSize | infnty | if e1 then e2
        | if e1 then e2 else e3
        | let x : τ = e1 in e2 | ⊞ [e | u ← g]
        | e; e | u.a | x = e.
        | for(u : g){e} | send(u, x) | halt

```

Figure 3: The syntax of Δ . Language forms that are highlighted are forms that are not visible to the user, and that are part of the internal representation of the AST used by the compiler.

the vertex in order to allow us to send incremental messages—or Δ -messages—and the definition of the Δ -messages that we send are defined in terms of (or at least need to take into account) one another.

5 THE LANGUAGE

The syntax of the language that we are compiling from is a small imperative query language that we call Δ , and is defined in Figure 3. It is designed to be small and not particularly special in any way since we want the transformations we express over it to be easily transferable to other languages. Further, we want the compilation process to be as uncomplicated as possible for expository purposes.

A Δ program consists of an initialization expression $\text{init}\{e\}$ that is run at each vertex at the start of the computation, and before any communication has taken place. After the initialization expression has been run, each statement is run sequentially, in-order, and homogeneously across all vertices. Each statement can either be a single $\text{step}\{e\}$ which represents a computation to be run for one single superstep, or an $\text{iter } i \{e_1\} \text{ until } \{e_2\}$ which keeps iterating e_1 until the expression e_2 is satisfied.

Being a small imperative language, the expressions in Δ are largely straightforward, with the only two interesting aspects being the $\boxplus[e|u \leftarrow g]$ form which functions as an aggregation operation over the vertices represented by g , and the vertex-field access form $u.a$ which states what field of the other vertex u we are accessing within the aggregation expression e . Thus if our vertex-state had a field amt , and we wanted to compute the sum of the amts of our neighboring vertices we could do so with the following aggregation operation

$$+ [u.\text{amt} | u \leftarrow \#neighbors] \quad (2)$$

Where we *always* aggregate over a graph-expression g with $\#in$ and $\#out$ corresponding to in- and out-vertices in the case that the graph is directed, and $\#neighbors$ referring to the neighboring vertices in the case that the graph is undirected.⁶

This aggregation operation represents a useful—but not crucial for our optimization techniques—difference from the normal programming model exposed for vertex-centric computation: the vast majority are *push-based* where the programmer directly sends and receives messages,⁷ whereas in Δ the language is *pull-based* where the (user-visible) vertex function directly operates over the value of its neighbors, as opposed to the messages sent to it by other vertices e.g., $\#in$, or $\#out$ neighbors.⁸

The difference between the push-based and pull-based programming models can be seen by comparing the push-based version of PageRank that we saw earlier in Figure 1 where the programmer directly sends messages on one superstep that will be received as messages on the next superstep, and the following definition of PageRank in Δ

```

init {
  // initialize the values
  local v1 : float = 1 / graphSize;
  local pr : float = v1 / |#neighbors|
};
iter i {
  // sum neighbors PageRanks
  let sum : float = + [ u.pr | u <- #neighbors ] in
  // calculate new value and new pagerank for
  // neighbors to see next superstep
  v1 = 0.15 + 0.85 * (sum / graphSize);
  pr = v1 / |#neighbors|
}

```

⁶In the case that the graph is undirected, $\#in$ and $\#out$ mean the same thing as $\#neighbors$.

⁷And in fact, all underlying implementations are push-based.

⁸Although in the compilation process we transform every “pull” into an appropriate “push” from the other vertices.

```

} until {
  i >= 30 // stop after 30 iterations
}

```

As we will see when we start compiling the language, specifically demarcating aggregation operations in this manner will make it easier for us to determine the best places to incrementalize the program. However, having these forms is not crucial to the optimizations that we present, and the aggregation points could be determined through a fairly straightforward flow-analysis on the program.

6 IMPLEMENTATION

Notation. Since the transformations that we will be performing on the AST are traversal independent, we use a context-based rewriting notation to detail the program transformations that we perform. Since they only take place over expressions, the contexts can be viewed as expressions with “holes”. Thus when we write $C[]$ we will mean an expression with a hole in it such as $C[] = \text{if } e_1 \text{ then } [] \text{ else } e_2$, and $C[e_3]$ represents a filling of this hole with the expression e_3 ; $C[e_3] = \text{if } e_1 \text{ then } e_3 \text{ else } e_2$.

The holes in our contexts are restricted to only appearing in places where an expression would be (i.e., expression context). Thus the following context would be *invalid* $C[] = u.[]$ since the hole does not appear in an expression context. We say that $C[e_1] = e$ if there exists a valid context $C[]$ such that $C[e_1] = e$, and we will also say in this case that $e_1 \in e$. Using this notation, we define the context-based compilation of an expression e by saying that $C[e_1] \rightsquigarrow C[e'_1]$ means that for every $e_1 \in e$, e_1 is transformed into e'_1 . Thus after this $e_1 \notin e$.

Before moving to the description of our transformations, we recall that we distinguish variables that are fields in our vertex state by using an underlined teletype font for them. Furthermore, before all of the below passes are run a type-annotation pass has been run on our source program, so we can always get the type of any expression e by calling the $\text{typeOf}(e)$ function.

6.1 Aggregation Conversion

In this pass we convert communication that is based on the pull-based model of $\mathbf{A}$ to the push-based model that we need for Pregel. In order to do this we employ the same technique as other pull-based vertex-centric languages such as Palgol [20], Green-Marl [8], and Fregel [5], where the key realization is that since the computation is homogeneous across the graph a “pull” on a vertex, is a “push” by its neighbors on the previous superstep. Since this is a standard compilation technique we don’t go into the details here and simply summarize the process as follows:

- We first convert all aggregation expressions that are not the immediate right hand side of an assignment (or let) and bind the result of the aggregation to a new variable and then substitute this variable in for occurrences of that aggregation.⁹ After this *all* aggregations that we will encounter are of the form $x = \boxplus[e \mid u \leftarrow g]$.¹⁰
- At the first superstep (i.e. immediately after initialization of the vertex state) send the data from the neighbors perspective: field accesses of a vertex within an aggregation expression

such as $u.\underline{\text{amt}}$ in Equation (2) become sends of that field ($\underline{\text{amt}}$) to that vertex’s neighbors.

- At subsequent steps every vertex inspects the message list (messages in Figure 1) that it receives in order to obtain values for its neighbors, and then executes the aggregation, so the following transformation is performed:

$$C[x = +[u.\underline{\text{amt}} \mid u \leftarrow g]] \rightsquigarrow C \left[\begin{array}{l} \text{let } tmp : \tau = \text{default_init}(+, \tau) \text{ in} \\ \text{for } (m : \text{messages}) \{ tmp = tmp + m.\underline{\text{amt}}; \}; \\ x = tmp \end{array} \right] \quad (3)$$

where τ is the type of $\underline{\text{amt}}$, and $\text{default_init}(\boxplus, \tau)$ returns the default initialization for the type τ and aggregation operation $\boxplus$. For the above example, since the aggregator is $+$, we get that $\text{default_init}(+, \tau) = 0$.¹¹

6.2 Adding Vertex State

In order to to remember the previous message sent from one superstep to the next, for each expression e that appears within a $\text{send}(u, e)$ expression we add a field to the vertex state to store the computed value for e , unless e is already a field of the vertex. This can be viewed as a type of A-normalization [15] except instead of binding to let-bound variables we are binding to fields in our vertex state:

$$C[\text{send}(u, e)] \rightsquigarrow C \left[\begin{array}{l} \text{freshVar}_e = e; \\ \text{send}(u, \underline{\text{freshVar}}_e) \end{array} \right] \quad (4)$$

where freshVar_e is a unique variable name and is associated with the type $\tau_e = \text{typeOf}(e)$. We then add, or “bind” each of these new fields to our vertex state:

```

VertexState {
  ...
   $\tau_e$   $\underline{\text{freshVar}}_e$ ;
}

```

We then remember each of these fields of our vertex that is accessed within sends for the next pass.

6.3 Inserting Change Checks

In the previous pass we made sure that we only send fields that are unmodified from our vertex state, and we kept track of all of the fields that were ever sent. We’ll call these the *externally visible fields* of the vertex, since they are the only things that other vertices can know about any other vertices.

We now add a “dirty-bit” to the vertex state to track whether any of these externally-visible states have changed during a computational step. Since we haven’t sent anything at the first superstep, the dirty-bit is pre-set in the initial vertex state:

```

VertexState {
  ...
  bool  $\underline{\text{dirtied}} = \text{true}$ ;
}

```

At the beginning of each superstep (except the first) we save the current state of each externally visible field x_f in a temporary variable o_f . Then whenever we encounter an assignment to that

⁹i.e., we A-normalize [15] with respect to aggregations.

¹⁰Or let bound, but the reasoning for these is identical so we will WOLOG only handle assignments in this section.

¹¹It is important to pass in the type along with the aggregation operation to default_init since for certain aggregation operations we would need to return the maximum or minimum value for τ .

field we also possibly dirty this dirty-bit by comparing the new value to our saved value of that field at the beginning of the superstep.

$$C[\underline{x_f} = e] \rightsquigarrow C \left[\begin{array}{l} \underline{x_f} = e; \\ \underline{dirtied} = \underline{dirtied} \parallel (\underline{x_f} \neq o_f) \end{array} \right] \quad (5)$$

Now, when we encounter sends of a field to neighboring vertices, we check whether the dirty-bit has been set

$$C[\text{send}(u, \underline{x_f})] \rightsquigarrow C \left[\begin{array}{l} \text{if}(\underline{dirtied}) \\ \text{then send}(u, \underline{x_f}) \end{array} \right] \quad (6)$$

Furthermore, when we encounter such a send in a broadcast setting we can lift the if-expression outside of the loop and prevent execution of the loop entirely if the message to be broadcast has not changed:

$$C[\text{for}(u : g)\{\text{send}(u, \underline{x_f})\}] \rightsquigarrow C \left[\begin{array}{l} \text{if}(\underline{dirtied}) \text{ then} \\ \text{for}(u : g)\{ \\ \text{send}(u, \underline{x_f}) \\ \} \end{array} \right] \quad (7)$$

6.4 Incrementalizing Aggregations

Having altered the program so messages are sent only when the value of has changed from the most recently sent value, we turn to incrementalizing the interaction between the vertices in the graph computation as a whole. To do this, we first alter aggregations while also taking into account how the $\Delta_u^m(m')$ function in the next pass will be determined.

We alter the aggregations by hoisting each aggregation operation so that it is memoized in the vertex state. We thus go back and change the operation that was performed in Equation (3), and instead perform the following transformation:

$$C[x = \boxplus[e | u \leftarrow g]] \rightsquigarrow C \left[\begin{array}{l} \text{for}(m : \text{messages})\{ \\ \quad \underline{\text{aggAccum}} = \underline{\text{aggAccum}} \boxplus e \\ \}; \\ x = \underline{\text{aggAccum}} \end{array} \right] \quad (8)$$

where $\underline{\text{aggAccum}}$ is a unique variable name associated with that particular aggregation expression. We then add this accumulator to the vertex state so that the previous value is remembered from one superstep to the next.

```
VertexState {
  ...
   $\tau$   $\underline{\text{aggAccum}}$  = default_init( $\boxplus$ ,  $\tau$ );
}
```

where τ and default_init are just as in Section 6.1. Thus this can be seen as performing the same function as the transformation that was done in Equation (3) with the exception that instead of let-binding to a temporary variable as is standard when converting to A-normal form, we instead use a vertex field so that we can use this to memoize the aggregations. Thus each aggregation within the vertex function starts with the value from the previous superstep. This then means that the messages that are sent only need to say how to change the previous value to get the new one, or in other words how to calculate the delta from the previous value.

6.4.1 Dealing with multiplicative operations. Memoizing things in the above manner works well in most cases, however in multiplicative contexts, such as with $\times$ or $\&\&$ aggregation operators, if we receive a “nullary-message” of 0 or false our accumulator is now nulled out, and there is no way to recover to a correct non-null accumulator while still keeping a meaningful-messages only policy—even if the offending nullary-message subsequently becomes a non-nullary value. In order to solve this problem, if we encounter a multiplicative aggregation operation instead of tracking one value we in fact track three values locally within the vertex’s state. These consist of the non-nulled result ($\underline{\text{nnAcc}}$), the number of nullary expressions ($\underline{\text{aggNulls}}$) not included in the $\underline{\text{nnAcc}}$ field, and a final value for the aggregation at that superstep ($\underline{\text{aggAccum}}$). When we have no nullary expressions in the aggregation (i.e., $\underline{\text{aggNulls}} == 0$) we set the value of $\underline{\text{aggAccum}}$ to the value of the non-null result $\underline{\text{nnAcc}}$, however if we still have nullary expressions in the aggregation ($\underline{\text{aggNulls}} > 0$) then we set the final value of $\underline{\text{aggAccum}}$ to the nullary value for the type and operation of the aggregation. Thus when we receive a nullary-message we increment the number of nullary expressions in our aggregation but we do not do anything else.¹² On the other hand, when we receive a non-nullary message, we check the tag of the message to see if it was a previously null message, and decrement the number of nullary elements in the aggregation ($\underline{\text{aggNulls}}$) if so. We thus perform the following transformation for multiplicative aggregations

$$C[x = \boxplus[e | u \leftarrow g]] \rightsquigarrow C \left[\begin{array}{l} \text{for}(m : \text{messages})\{ \\ \quad \text{if}(\text{is_nullary}(m)) \\ \quad \text{then } \underline{\text{aggNulls}} = \underline{\text{aggNulls}} + 1; \\ \quad \text{else} \\ \quad \quad \underline{\text{nnAcc}} = \underline{\text{nnAcc}} \boxplus e; \\ \quad \quad \text{if}(\text{prev_nullary}(m)) \\ \quad \quad \text{then } \underline{\text{aggNulls}} = \underline{\text{aggNulls}} - 1; \\ \quad \} \\ \quad \text{if}(\underline{\text{aggNulls}} == 0) \\ \quad \text{then } \underline{\text{aggAccum}} = \underline{\text{nnAcc}}; \\ \quad \text{else } \underline{\text{aggAccum}} = \text{nullary_elem}(\boxplus, \tau); \\ \quad x = \underline{\text{aggAccum}} \end{array} \right] \quad (9)$$

and add the following fields to the vertex state

```
VertexState {
  ...
  int  $\underline{\text{aggNulls}}$  = 0;
   $\tau$   $\underline{\text{nnAcc}}$  = default_init( $\boxplus$ ,  $\tau$ );
   $\tau$   $\underline{\text{aggAccum}}$  = default_init( $\boxplus$ ,  $\tau$ );
}
```

The $\Delta_u^m(m')$ needs to track nullary messages values and tag non-nullary messages with whether or not the previous message (m) from that vertex was a nullary message or not.

6.5 Δ -Message Insertion

Now that we have memoized the aggregation operations between supersteps, we now turn to we determine the $\Delta_u^m(m')$ function that computes the incrementalized messages. However before we do

¹²Due to the meaningful-messages only policy, we know that it was a non-previously null message, and we therefore increment the number of nullary results.

this, we insert calls to calculate Δ -messages. Since we already saved the values of our externally visible fields by the transformation in Section 6.3 we use this saved value when computing the Δ -message. We thus perform the following transformation

$$C[\text{send}(u, x_f)] \rightsquigarrow C[\text{send}(u, \Delta_u^{of}(x_f))] \quad (10)$$

Now that calls to our message incrementalization function have been inserted, we determine what exactly this function should be by the following equation:

$$x \boxplus m' \simeq (x \boxplus m) \boxplus \Delta_u^m(m') \quad (11)$$

where x is an accumulator, $\boxplus$ is an aggregation operation, and $\simeq$ is (as before) denotational equality. In more intuitive language, for each aggregation we track the fields that are sent to that aggregation, and the aggregation operation performed (e.g., for the PageRank example we would get $\{+ \mapsto \text{pr}\}$). Then for an aggregation operation $\boxplus$, we synthesize $\Delta_u^m(m')$ such that Equation (11) holds.

In the case that the aggregation is multiplicative (subject to possible nulling) we tag each message with whether or not the previous message was nullary or not. Tagging for multiplicative operations is very useful here, since e.g., in the $\boxplus = \times$ case, we need to avoid division by zero in the $\Delta_u^m(m')$ function, and in this case we calculate the incrementalized message by the following:

$$\Delta_u^m(m') = \begin{cases} m'/m & \text{if } m \neq 0 \\ \text{tag}(m') & \text{if } m = 0 \end{cases}$$

6.6 Addition of Halts

While this pass is not necessary for correctness, the transformations that we have performed in order to ensure a meaningful-only messaging policy exposes a powerful characteristic that we can take advantage of. Recall that in our meaningful-only messaging policy, messages also serve as cache invalidations, thus once a vertex has computed a specific value and sent its messages, the only way the values of the messages that it sends can change is by it receiving new messages.¹³ However any *halted* vertex will be woken up if it receives any new messages. Since we only send meaningful messages, this means that after sending its messages, *every vertex can halt after every superstep*.

As opposed to the previous transformations which all applied to expressions e in Δ , this pass is instead a straightforward transformation on statements in Δ which simply inserts *halts* at the end of every statement given by the following:

$$\begin{aligned} \text{step}\{e\} &\rightsquigarrow \text{step}\{e; \text{halt}\} \\ \text{iter } i \{e\} \text{ until } \{e\} &\rightsquigarrow \text{iter } i \{e; \text{halt}\} \text{ until } \{e; \text{halt}\} \end{aligned} \quad (12)$$

7 EVALUATION

In this section we evaluate the usefulness of the transformations that we have presented by evaluating them on a number of standard benchmarks. In particular we compare the number of messages sent, and total computation time for PageRank (PG) [14], Single-Source-Shortest Paths (SSSP), Connected Components (CC), and a non-converging version of Hyperlink-Induced Topic Search (HITS) [10] where we perform the hub and authority updates

¹³We assume here that all operations are deterministic, and we don't have access to functions such as `rand()`.

Table 1: Datasets used in our evaluation.

Dataset	Type	V	E
Wikipedia	Directed	18.27M	136.54M
LiveJournal-DG	Directed	4.85M	68.48M
Facebook	Undirected	59.22M	185.04M
LiveJournal-UG	Undirected	3.99M	34.68M

Table 2: Size (in bytes) of vertex state for Δ , and Δ without incrementalization (Δ^*) as compared to Pregel+ and Palgol vertex state.

Variant	PageRank	SSSP	CC	HITS
Δ	48 B	48 B	48 B	80 B
Δ^*	40 B	40 B	40 B	64 B
Palgol	40 B	64 B	40 B	64 B
Pregel+	32 B	40 B	32 B	56 B

simultaneously. For each of these benchmarks, we compare Δ to reference implementations for each benchmark written in Pregel+ [19], and compare Δ to itself *without* message-reduction optimizations (Δ^*). All experiments were conducted on an Amazon EC2 cluster with 8 m4.xlarge nodes, each having 4 vCPUs and 16GB of memory running two workers per node, and connected by 750Mbps ethernet. The datasets that we use are listed in Table 1.

7.1 Change in Vertex State Size

Since we want the compiled and incrementalized program to use as little additional memory as possible as compared to a non-incrementalized version, we compare the size of the vertex states generated by both Δ and Δ^* , and compare this to the vertex states generated by Palgol [20]—another language that compiles to Pregel+—in addition to the size of the vertex state in the reference implementations in Pregel+.

As we see from Table 2 compiled vertex state sizes in general are larger than their comparable vertex state in a hand-written Pregel+ program. This difference is due to the fact that the source program needs to be compiled to a state machine within the Pregel+ `compute()` function, and therefore additional internal state is required in order to track these states and transitions. However, we are not really interested in the size difference between the compiled and hand-written versions. Instead the more useful comparison is between Δ and the two non-incrementalized languages Δ^* and Palgol. In this case we see that while incrementalization adds some additional space overhead, this additional space is fairly minimal compared to the overall size of the vertex state.

7.2 Performance and Message Reduction

In order to evaluate the viability of the transformations that we have presented, we want to show that we don't slow computations down in which the transformations are not applicable, and that the language itself is efficient enough to be useable, and therefore that these optimizations can be practical in the real-world. In order to evaluate both of these properties we compare Δ to itself without the optimizations that we have presented— Δ^* —along with the

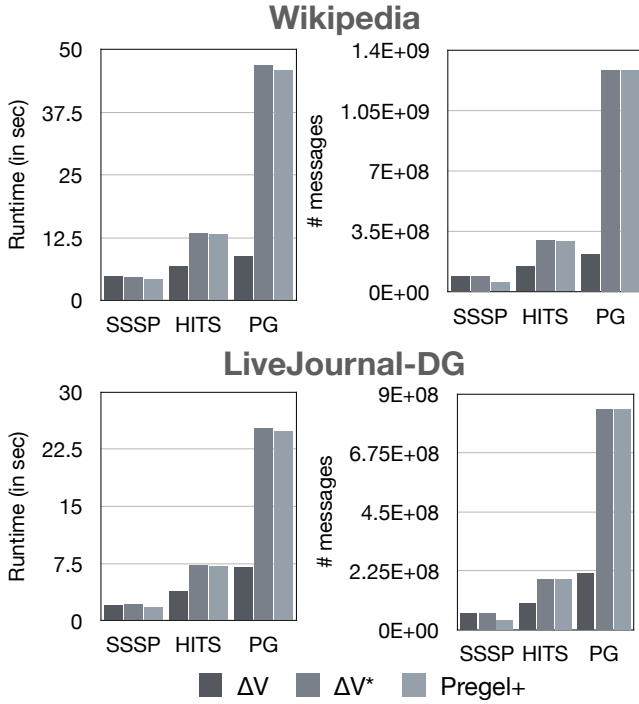


Figure 4: Execution time (left) and number of messages sent (right) for PageRank, SSSP, and HITS.

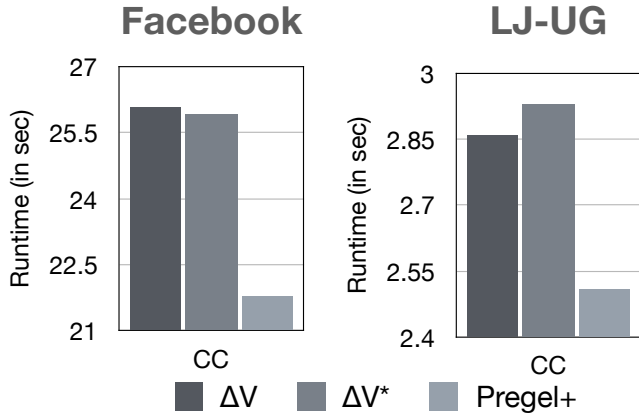


Figure 5: Execution time for Connected Components on Facebook (left) and LiveJournal-undirected (right).

same algorithms implemented in Pregel+. The results of these benchmarks are presented in Figures 4 and 5 and are obtained by taking the average of three runs. PageRank was run for 30 iterations, and HITS for 7 (5 after 2 initialization steps).

As we see from Figure 4, while Pregel+ is always faster than ΔV^* , on certain iterative algorithms (PG & HITS) the optimizations that we have performed in ΔV can increase performance significantly. With PageRank having an average speedup of 4.4X across the Wikipedia and LiveJournal graphs and HITS showing an average speedup of 1.9X as well, as compared to Pregel+. Furthermore, we see a 5.8X and 1.9X reduction in the number of messages sent for

the PageRank and HITS programs respectively.

On other programs that are less amenable to incrementalization, such as CC (Figure 5) and SSSP (Figure 4), we do not see any performance improvement at all. With ΔV^* and ΔV in fact sending the exact same number of messages in both cases¹⁴. The lack of a speedup from our optimizations for these programs makes sense; in both CC and SSSP the standard algorithm is in a sense “pre-incrementalized”: the compute() function only sends new messages when a new value has been attained at that vertex. The most important thing then for these two benchmarks is that the optimizations that we have performed in ΔV have *not slowed the computation down*. Thus, while our optimizations have the ability to drastically increase the performance of certain types of programs, they do not negatively effect programs when they have nothing to contribute.

8 RELATED WORK

The majority of previous work for vertex-centric computation has taken a library-oriented, or shallow embedding view such as Pregel [12], Pregel+ [19], Apache Giraph [6], Graphlab [11] and many others [2, 9, 16]. While this library-based approach has the benefit of being lightweight and portable, a library-based system is unable to gather information about the program as a whole, or transform it. Thus a library-oriented approach is not suitable for determining and preventing non-meaningful messages, and we need to instead compile, or deeply-embed our language.

For compilation, we have followed the approach of other deeply-embedded, or language-based approaches to vertex-centric programming such as Palsgol [20], Green-Marl [8], and Fregel [5] in which the programmer writes a vertex-centric algorithm kernel which is then compiled to a vertex-centric (library-based) framework. However, the approach that we have taken towards incrementalization of the language can be seen as an extension of different techniques in the wider library-based vertex-centric community over the past years. Particularly, Apache Flink [2] has supported delta iterations—which are very similar to our idea of an incrementalization of a vertex-centric program—for a number of years, however the user still needs to hand-write and transform their code in order to use these capabilities. Other compiled DSLs for vertex-centric computations have also worked on communication optimizations as well [8, 20], however none have sought to automatically incrementalize the program during the compilation process for vertex-centric computations, or to reduce the number of messages sent in the compiled program in this manner.

9 CONCLUSIONS & FUTURE WORK

In this paper we have introduced a family of program transformations that together incrementalize the interaction between vertex computations and through this incrementalize the computation as a whole (Section 6), and demonstrated the usefulness of these transformations (Section 7). We have shown that programs written in our prototype language ΔV that have been incrementalized by these transformations always achieve *at least* the same performance as programs that have not been optimized, and that these optimizations have the ability to drastically increase performance on iterative

¹⁴For CC ΔV , ΔV^* , and Pregel+ sent the exact same number of messages so this graph has been elided.

algorithms: we can therefore perform these transformations on a source program and not worry about possibly slowing the computation down. Furthermore, these transformations are straightforward, readily applicable to other (compiled) DSLs for vertex-centric computation, and not only do not restrict other standard optimization techniques for vertex-centric programs but open up other low-level optimization techniques. Indeed, while we have demonstrated the effectiveness of incrementalization there are still a number of future directions and possibilities to explore that are directly related to the method of automatic incrementalization that we have presented.

Future Work

One future direction for this work is to extend the incrementalization policy to allow removing vertices from the graph during the computation; while $\mathcal{N}$ is a query language and thus cannot manipulate the graph, the incrementalization approach that we have presented can be extended to handle these manipulations. A vertex could be deleted from the graph in the same manner as described in Pregel [12], however with the addition that the vertex being deleted first broadcasts a message that zeros out the value of the vertex to its neighbors before the deletion of the edges is performed (*i.e.* the vertex needs to set the neighbors value for that vertex's most recent message to a unitary value.).

Recent work in vertex-centric frameworks has explored separating messages into different “channels” based upon how those messages are meant to be used in the computation, and the overarching communication patterns known about the computation. Once messages are separated in this manner specific messaging policies such as request-respond, broadcast, and others can be used to optimize the different types of communication that can happen between vertices. It would be interesting to explore how further optimization of the compiled program could be done by performing static analysis during compilation to determine what types of message-passing channels should be used, and to possibly build a channel type for incremental messages.

Another future direction is to allow the programmer to define an “allowable slop” parameter ϵ to the program where a message value is counted as changed if it is no longer within ϵ of the old message. Thus we wouldn't send a message from a vertex as long as the new message is within the ϵ of the most recently sent message, but the change to the outgoing message value is tracked internally within the vertex's state from one superstep to another. Once the message value differs by at least ϵ from the most recently sent message it is counted as changed, and a message is sent. With this view, the work described in this paper could be seen as a degenerate case of this where $\epsilon = 0$.

We currently do not take advantage of the halt-by-default policy that we have in a compiled $\mathcal{N}$ program. However, we believe that there is significant performance benefits that can be seen by changing the way Pregel+ determines (and schedules) which vertices should be run at any given superstep under this policy. In a none halt-by-default setting, at each superstep *each vertex in the graph* needs to be accessed in order to determine if it should be run (since the vertex may be active). However, in a halt-by-default setting, with the exception of the first superstep, the vertices that run at any given step are determined by the messages that are received at the previous superstep. Using this knowledge, we could use the message passing framework to build a work-queue based

scheduler for the computation, as opposed to how it currently stands where we need to loop through each and every vertex in the graph at every superstep, in order to determine which vertices should run.

REFERENCES

- [1] Yufei Cai, Paolo G. Giarrusso, Tillmann Rendel, and Klaus Ostermann. 2014. A Theory of Changes for Higher-order Languages: Incrementalizing Lambda-calculi by Static Differentiation. In *Proceedings of the 35th ACM SIGPLAN Conference on Programming Language Design and Implementation (PLDI '14)*. ACM, New York, NY, USA, 145–155. DOI: <http://dx.doi.org/10.1145/2594291.2594304>
- [2] Paris Carbone, Asterios Katsifodimos, Stephan Ewen, Volker Markl, Seif Haridi, and Kostas Tzoumas. 2015. Apache flink: Stream and batch processing in a single engine. *Bulletin of the IEEE Computer Society Technical Committee on Data Engineering* 36, 4 (2015).
- [3] Avery Ching, Sergey Edunov, Maja Kabiljo, Dionysios Logothetis, and Sambavi Muthukrishnan. 2015. One trillion edges: Graph processing at facebook-scale. *Proceedings of the VLDB Endowment* 8, 12 (2015), 1804–1815.
- [4] Ravikant Dindokar, Neel Choudhury, and Yogesh Simmhan. 2016. A meta-graph approach to analyze subgraph-centric distributed programming models. In *Big Data (Big Data), 2016 IEEE International Conference on*. IEEE, 37–47.
- [5] Kento Emoto, Kiminori Matsuzaki, Zhenjiang Hu, Akimasa Morihata, and Hideya Iwasaki. 2016. Think Like a Vertex, Behave Like a Function! A Functional DSL for Vertex-centric Big Graph Processing. In *Proceedings of the 21st ACM SIGPLAN International Conference on Functional Programming (ICFP 2016)*. ACM, New York, NY, USA, 200–213. DOI: <http://dx.doi.org/10.1145/2951913.2951938>
- [6] Apache Software Foundation. 2012. Giraph. <http://giraph.apache.org/>. (2012).
- [7] Alessio Guerrieri, Alberto Montresor, and Simone Centellegher. 2016. ETSCH: Partition-centric Graph Processing. In *Ubiquitous Intelligence & Computing, Advanced and Trusted Computing, Scalable Computing and Communications, Cloud and Big Data Computing, Internet of People, and Smart World Congress (UIC/ATC/SCCom/CBDCCom/IoP/SmartWorld), 2016 Intl IEEE Conferences*. IEEE, 706–713.
- [8] Sungpack Hong, Hassan Chafi, Edic Sedlar, and Kunle Olukotun. 2012. Green-Marl: A DSL for Easy and Efficient Graph Analysis. In *Proceedings of the Seventeenth International Conference on Architectural Support for Programming Languages and Operating Systems (ASPLOS XVII)*. ACM, New York, NY, USA, 349–362. DOI: <http://dx.doi.org/10.1145/2150976.2151013>
- [9] Zuhair Khayyat, Karim Awara, Amani Alonazi, Hani Jamjoom, Dan Williams, and Panos Kalnis. 2013. Mizan: a system for dynamic load balancing in large-scale graph processing. In *Proceedings of the 8th ACM European Conference on Computer Systems*. ACM, 169–182.
- [10] Jon M. Kleinberg. 1999. Hubs, Authorities, and Communities. *ACM Comput. Surv.* 31, 4es, Article 5 (Dec. 1999). DOI: <http://dx.doi.org/10.1145/345966.345982>
- [11] Yucheng Low, Danny Bickson, Joseph Gonzalez, Carlos Guestrin, Aapo Kyrola, and Joseph M Hellerstein. 2012. Distributed GraphLab: a framework for machine learning and data mining in the cloud. *Proceedings of the VLDB Endowment* 5, 8 (2012), 716–727.
- [12] Grzegorz Malewicz, Matthew H Austern, Aart JC Bik, James C Dehnert, Ilan Horn, Naty Leiser, and Grzegorz Czajkowski. 2010. Pregel: a system for large-scale graph processing. In *Proceedings of the 2010 ACM SIGMOD International Conference on Management of data*. ACM, 135–146.
- [13] Robert Ryan McCune, Tim Weninger, and Greg Madey. 2015. Thinking like a vertex: a survey of vertex-centric frameworks for large-scale distributed graph processing. *ACM Computing Surveys (CSUR)* 48, 2 (2015), 25.
- [14] Lawrence Page, Sergey Brin, Rajeev Motwani, and Terry Winograd. 1999. *The PageRank citation ranking: Bringing order to the web*. Technical Report. Stanford InfoLab.
- [15] Amr Sabry and Matthias Felleisen. 1992. Reasoning About Programs in Continuation-passing Style.. In *Proceedings of the 1992 ACM Conference on LISP and Functional Programming (LFP '92)*. ACM, New York, NY, USA, 288–298. DOI: <http://dx.doi.org/10.1145/141471.141563>
- [16] Semih Salihoglu and Jennifer Widom. 2013. GPS: A Graph Processing System. In *Proceedings of the 25th International Conference on Scientific and Statistical Database Management (SSDBM)*. ACM, New York, NY, USA, Article 22, 12 pages. DOI: <http://dx.doi.org/10.1145/2484838.2484843>
- [17] Yuanyuan Tian, Andrey Balmin, Severin Andreas Corsten, Shirish Tatikonda, and John McPherson. 2013. From think like a vertex to think like a graph. *Proceedings of the VLDB Endowment* 7, 3 (2013), 193–204.
- [18] Leslie G Valiant. 1990. A bridging model for parallel computation. *Commun. ACM* 33, 8 (1990), 103–111.
- [19] Da Yan, James Cheng, Yi Lu, and Wilfred Ng. 2015. Effective techniques for message reduction and load balancing in distributed graph computation. In *Proceedings of the 24th International Conference on World Wide Web*. International World Wide Web Conferences Steering Committee, 1307–1317.
- [20] Yongzhe Zhang, Hsiang-Shang Ko, and Zhenjiang Hu. 2017. Palgol: A High-Level DSL for Vertex-Centric Graph Processing with Remote Data Access. In *Programming Languages and Systems*, Bor-Yuh Evan Chang (Ed.). Springer International Publishing, Cham, 301–320.