

SuSe: Summary Selection for Regular Expression Subsequence Aggregation over Streams

STEVEN PURTZEL, Humboldt-Universität zu Berlin, Germany

MATTHIAS WEIDLICH, Humboldt-Universität zu Berlin, Germany

Regular expressions (Regex) are an essential tool for pattern matching over streaming data, e.g., in network and security applications. The evaluation of Regex queries becomes challenging, though, once subsequences are incorporated, i.e., characters in a sequence may be skipped during matching. Since the number of subsequence matches may grow exponentially in the input length, existing Regex engines fall short in finding *all* subsequence matches, especially for queries including Kleene closure.

In this paper, we argue that common applications for Regex queries over streams do not require the enumeration of all *distinct* matches at *any* point in time. Rather, only an aggregate over the matches is typically fetched at specific, yet unknown time points. To cater for these scenarios, we present SuSe, a novel architecture for Regex evaluation that is based on a query-specific summary of the stream. It employs a novel data structure, coined STATESUMMARY, to capture aggregated information about subsequence matches. This structure is maintained by a summary selector, which aims at choosing the stream projections that minimize the loss in the aggregation result over time. Experiments on real-world and synthetic data demonstrate that SuSe is both effective and efficient, with the aggregates being based on several orders of magnitude more matches compared to baseline techniques.

CCS Concepts: • **Information systems** → **Data management systems**; **Data streams**; **Stream management**.

Additional Key Words and Phrases: Regular Expression, Stream Processing, Stream Summary

ACM Reference Format:

Steven Purtzel and Matthias Weidlich. 2025. SuSe: Summary Selection for Regular Expression Subsequence Aggregation over Streams. *Proc. ACM Manag. Data* 3, 3 (SIGMOD), Article 222 (June 2025), 27 pages. <https://doi.org/10.1145/3725359>

1 Introduction

The evaluation of regular expressions (Regex) over streaming data is at the core of applications in various domains, such as network monitoring [31, 38, 41, 42], financial fraud detection [54, 57], or infrastructure security [12, 37]. Regex have been studied extensively in theoretical computer science [30, 33] and automata-based approaches for their evaluation, most prominently the Thompson construction [55] and its derivatives, have been around for decades.

However, existing Regex engines, such as PCRE2 [18], pcregrep [28], Oniguruma [2], Boost.Regex [39], RE2 [50], and TRE [1] are limited in their ability to consider subsequences during evaluation: First, they lack native support for skipping characters in a sequence when constructing matches. This is not a matter of expressiveness, though, since the skipping of characters can be encoded explicitly, e.g., by transforming a regular expression $\gamma = abc$ into $\gamma' = \Sigma^*a\Sigma^*b\Sigma^*c\Sigma^*$ with Σ being the underlying alphabet. However, such explicit encoding significantly increases the per-match state and

Authors' Contact Information: Steven Purtzel, Humboldt-Universität zu Berlin, Berlin, Germany, purtzesc@hu-berlin.de; Matthias Weidlich, Humboldt-Universität zu Berlin, Berlin, Germany, matthias.weidlich@hu-berlin.de.



This work is licensed under a Creative Commons Attribution 4.0 International License.

© 2025 Copyright held by the owner/author(s).

ACM 2836-6573/2025/6-ART222

<https://doi.org/10.1145/3725359>

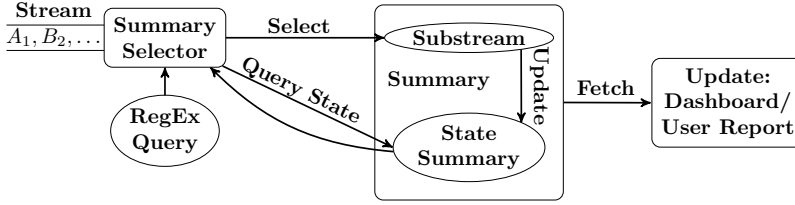


Fig. 1. Architecture for RegEx evaluation over streams.

changes the underlying language. Second, they do not support enumerating *all* matches, allowing only greedy or lazy matching that returns a single match. REMatch [51] is a notable exception in that it supports the enumeration of *all* matches; yet, it suffers from performance issues due to the exponential growth of subsequence matches with input size, as we demonstrate later empirically.

Complex Event Processing (CEP) engines like FlinkCEP [20] allow for subsequence matching using automata-based methods. They evaluate sequence queries grounded in regular expressions, enhanced with predicates and time window constraints. However, these engines also face performance issues caused by the need to manage exponentially many matches. The CORE engine [10] seeks to address this by representing matches in an efficient data structure. Still, it exhaustively enumerates matches, leading to significant performance bottlenecks, as demonstrated by our experiments.

In this paper, we argue that RegEx evaluation with subsequences may be optimized based on the following two observations:

- (1) Many application scenarios demand that all matches of a regular expression are incorporated, whereas the output is given in terms of an aggregate (e.g., count or sum) computed over all matches. Such scenarios are referred to as *event trend aggregation* and dedicated solutions that are based on incremental computation of the aggregate have been developed for it, e.g., GRETA [47] and its successors [43, 46, 48, 52].
- (2) The matches of a regular expression, and aggregates over them, are not necessarily needed at any point in time. Rather, many application scenarios involve an external trigger that determines when the result is relevant. For instance, users may load an analysis report, open a dashboard, or refresh a visualization at certain, generally unknown time points.

Combining these two observations, we present SuSE, a novel architecture for RegEx evaluation. As shown in Fig. 1, it incrementally maintains a query-specific summary of the stream that includes a substream of selected elements along with aggregated information about query matches. This summary provides an approximate answer, which is generally more accurate than existing evaluation strategies, and which can be fetched at any point in time. From a technical point of view, SuSE makes the following contributions:

- (1) We present a model for RegEx subsequence matching with aggregations that is based on a summary of a stream. Based thereon, we formulate the problem of constructing an optimal summary that minimizes the expected information loss.
- (2) We introduce STATESUMMARY, a novel data structure to capture aggregated information about subsequence matches. We elaborate on the operations to maintain this structure.
- (3) We present an approach for summary selection, which aims at constructing an optimal summary. It is based on a cost model to assess the potential benefit of stream elements.

Below, we first provide a motivating example (§2) and introduce preliminaries (§3), before giving a formal problem statement (§4). Then, we define the STATESUMMARY (§5), as the basis for summary selection (§6). We present evaluation results (§7), demonstrating that SuSE computes aggregates

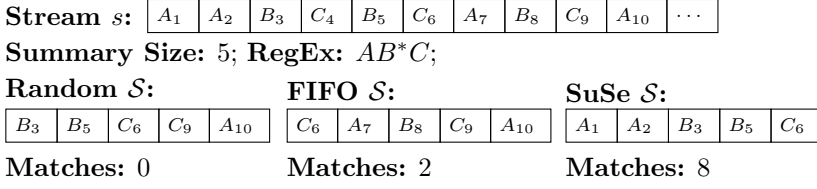


Fig. 2. Comparison of random, First-In-First-Out (FIFO), and our SuSe summary selection strategy.

based on several orders of magnitude more matches than baseline techniques. Finally, we review related work (§8) and conclude the paper (§9).

2 Motivating Example

Consider a bike rental service that aims to enhance operational efficiency by analyzing patterns within trip data streams, e.g., targeting patterns of events denoting short trips in busy areas. These patterns, indicative of user behavior, are defined using regular expressions. The inherent variability and complexity of real-world data streams, where a strictly contiguous occurrence of events within a pattern may be rare or too restrictive, motivates the search for subsequences, i.e., events in a pattern can occur in a non-contiguous order in the stream. Furthermore, streams are partitioned by attributes, e.g., bike ID or start location, to facilitate targeted analysis.

To inform operational decision making, a dashboard gives an overview of the current number of patterns, i.e., it derives a simple aggregate over the matches of a query that describes the relevant trips. Whenever the dashboard is refreshed, at unknown time points, the aggregate is updated and a subsequence of the stream that induced many pattern matches is made available for a quick plausibility assessment. Here, it is fine to trade a minor information loss due to approximate results for increased efficiency, as decision making is based on general trends rather than the exact results.

Fig. 2 showcases three strategies for retaining stream subsequences: Random, First-In-First-Out (FIFO), and SuSe. Given is a stream s , a summary size n , and a regular expression $\gamma = AB^*C$. Assuming that all events of s fall within the time window of the Regex query, we aim to select a subsequence S of s (i.e., an order-preserving selection of elements) no larger than n , that maximizes the count of subsequence matches upon evaluating the Regex against the summary S . For the stream in Fig. 2, we compare the strategies at the point in time after the event A_{10} was processed.

Random. A first baseline strategy randomly selects an element for eviction when appending a new one, once the summary size n is reached. This way, a subsequence was chosen which produces zero matches upon evaluating γ against S , wasting the entire space.

FIFO. The FIFO strategy always keeps the latest n elements in the summary S . Evaluating γ against S generates two matches, i.e., A_7C_9 and $A_7B_8C_9$. However, keeping C_6 in the summary, which cannot contribute to any matches, illustrates inefficient space usage.

SuSe. Using the SuSe framework presented in this work, we would arrive at a more effective stream summary resulting in eight matches: A_1C_6 ; A_2C_6 ; $A_1B_3C_6$; $A_2B_3C_6$; $A_1B_5C_6$; $A_2B_5C_6$; $A_1B_3B_5C_6$; and $A_2B_3B_5C_6$. Yet, the answer is still approximate (missing, for instance, matches with event B_8 of the stream).

3 Preliminaries

3.1 Regular Expressions

Languages. An *alphabet* $\Sigma = \{a_1, \dots, a_m\}$ is a set of *characters*. A *word* is a sequence of characters $w = \langle c_1, \dots, c_n \rangle$ of size $n = |w|$ over Σ , i.e., $c_i \in \Sigma$. The character at the i -th index in word w is denoted by $w[i]$, for $1 \leq i \leq |w|$. The concatenation of words $w = \langle c_1, \dots, c_n \rangle$ and $w' = \langle c'_1, \dots, c'_m \rangle$

is $w.w' = \langle c_1, \dots, c_n, c'_1, \dots, c'_m \rangle$. A *prefix* of w is any word of length $\leq n$ that starts from its first character.

A *subsequence* y of a word $w = \langle c_1, \dots, c_n \rangle$, denoted by $y \leq w$, is obtained by deleting zero or more characters from w without altering the order of the remaining characters. There are 2^n possible subsequences of w . Formally, a subsequence y is represented as $y = \langle c_{i_1}, \dots, c_{i_k} \rangle$, where $k \leq n$ and $1 \leq i_1 \leq \dots \leq i_k \leq n$.

To establish the relation between the indices in y and w , a *mapping function* $m: \{1, \dots, k\} \rightarrow \{1, \dots, n\}$ is defined such that $j \mapsto i_j$. This function maps the k indices of y to the corresponding indices in w , denoted as $y \leq_m w$. A *partial mapping* of a subsequence y' to w , denoted as $y' \leq_{pm} w$, involves any prefix y' of y that can be mapped to w using function pm . Note that multiple (partial) mapping functions may exist from a subsequence to a word.

EXAMPLE 1. Consider the alphabet $\Sigma = \{A, B, C, D\}$ and a word $w = \langle A_1, B_2, D_3, C_4, A_5, B_6, C_7, D_8 \rangle$. Here, $y = \langle A_1, B_2, D_3, C_4 \rangle$ is a subsequence with two different mappings: we have $y \leq_{m_1} w$ with $m_1: 1 \mapsto 1, 2 \mapsto 2, 3 \mapsto 3, 4 \mapsto 4$; and $y \leq_{m_2} w$ with $m_2: 1 \mapsto 1, 2 \mapsto 2, 3 \mapsto 3, 4 \mapsto 7$.

The set of all words of length $n \in \mathbb{N}_0$ over an alphabet Σ is Σ^n ; and the set of all possible words is $\Sigma^* = \bigcup_{n \geq 0} \Sigma^n$. A *language* L over Σ is any subset $L \subseteq \Sigma^*$, and the set of all prefixes of L is denoted as L_{pre} . The *empty word* is $\varepsilon = \Sigma^0$.

Let L_1 and L_2 be languages defined over an alphabet Σ . Then, we define concatenation, union, and Kleene star for languages as:

- $L_1 L_2 = \{x.y \mid x \in L_1, y \in L_2\}$,
- $L_1 \cup L_2 = \{x \in \Sigma^* \mid x \in L_1 \vee x \in L_2\}$, and
- $L_1^* = \bigcup_{n \geq 0} L_1^n$.

RegEx. We inductively define regular expression semantics. The symbols $\emptyset$, ε , and $a \in \Sigma$ are regular expressions, describing:

- the empty language $L(\emptyset) = \emptyset$;
- the language $L(\varepsilon) = \langle \varepsilon \rangle$, and
- for each character $a \in \Sigma$, the language $L(a) = \{\langle a \rangle\}$.

Let α and β be regular expressions, with languages $L(\alpha)$ and $L(\beta)$. Then, semantics of concatenation $\alpha\beta$, union $(\alpha|\beta)$, and Kleene star $(\alpha)^*$ are defined as follows: $L(\alpha\beta) = L(\alpha)L(\beta)$; $L(\alpha|\beta) = L(\alpha) \cup L(\beta)$; and $L((\alpha)^*) = L(\alpha)^*$.

3.2 Regular Expression Queries over Streams

Natural numbers represent time points, denoted as $\mathbb{T} \subseteq \mathbb{N}$. A *stream* $s = (\omega = \langle c_1, c_2, \dots \rangle, \lambda = \langle t_1, t_2, \dots \rangle, v = \langle v_1, v_2, \dots \rangle)$ comprises sequences of characters $c_i \in \Sigma$ (i.e., ω), timestamps $t_i \in \mathbb{T}$ (i.e., λ), and values $v_i \in \mathbb{N}$ (i.e., v). They induce stream elements $e_i = (\omega[i], \lambda[i], v[i])$ for $1 \leq i$, with $\omega[i]$ (or c_i), $\lambda[i]$ (or t_i), and $v[i]$ (or v_i) being the i -th elements of the respective sequences. When referencing a stream element e_i corresponding to a mapping function m , we use bracket notation, e.g., $\lambda[m(i)]$; if the index is clear, we refer to an element's attribute values by index, e.g., t_i , to reduce clutter.

The stream is assumed to be totally ordered by timestamps, i.e., $t_i < t_j$ for $1 \leq i < j$, which, in practice, may be achieved by techniques for out-of-order handling and tie-breaking mechanisms [45, 56]. The value v_i models a payload of e_i and could also be defined over a different domain, depending on the application scenario. Let $t \in \mathbb{T}$ be a time instance. A *substream* $s(t) = (\langle c_1, \dots, c_p \rangle, \langle t_1, \dots, t_p \rangle, \langle v_1, \dots, v_p \rangle)$ is the stream up to index p , containing only elements with timestamps $\leq t$.

A *regular expression query* $q = (\gamma, \tau, \mathcal{A})$ is a triple, comprising a regular expression γ (also q_γ), a time window $\tau \in \mathbb{T}$ (also q_τ), and an aggregate function $\mathcal{A}$ (also $q_{\mathcal{A}}$). A mapping $y \leq_m \omega$ of a word

$y \in L(\gamma)$ to ω is a *match* of a regular expression query q if the time window constraint is satisfied, i.e., $\lambda[m(k)] - \lambda[m(1)] \leq q_\tau$.

EXAMPLE 2. For $\Sigma = \{A, B, C\}$, consider the stream $s = (\omega, \lambda, \nu)$ with $\omega = \langle A_1, B_2, C_3, C_4 \rangle$ and $\lambda = \langle 1, 5, 6, 9 \rangle$. For the RegEx $q_\gamma = ABC$, two mappings between the only word $\langle A_1, B_2, C_3 \rangle$ in $L(q_\gamma)$ and ω exist: $m_1: 1 \mapsto 1, 2 \mapsto 2, 3 \mapsto 3$ and $m_2: 1 \mapsto 1, 2 \mapsto 2, 3 \mapsto 4$. Yet, with a time window $q_\tau = 5$, only m_1 is a match of the RegEx query.

A *partial match* occurs when a word $y \in L_{\text{pre}}(\gamma)$ maps to indices of ω and satisfying the time window constraint. It is a candidate for becoming a complete match as more characters of ω are examined. Notably, a complete match can also serve as a partial match. For instance, for $\gamma = A(B|C)D^*$, a match of the word $\langle A_1, B_2, D_3 \rangle \in L(\gamma)$ is also a partial match of $\langle A_1, B_2, D_3, D_4 \rangle \in L(\gamma)$.

Our work utilizes a streaming model of computation, in which a RegEx query $q = (\gamma, \tau, \mathcal{A})$ is received and preprocessed. We focus on the aggregate functions $\mathcal{A} \in \{\text{COUNT}, \text{SUM}, \text{AVG}\}$, which are defined over the matches of q . COUNT calculates the number of matches, SUM adds up the payload values of elements from all complete matches, and AVG is the ratio of SUM to COUNT.

3.3 Regular Expression Query Matching

Automata-Based Matching. Traditional matching of a regular expression q_γ against a stream (word) relies on a Non-deterministic Finite Automaton (NFA). An NFA is a tuple $\mathcal{N} = (Z, \Sigma, \delta, z_0, F)$, with Z as a finite set of states, Σ as the alphabet, $z_0 \in Z$ as the initial state, $\delta: Z \times \Sigma \rightarrow 2^Z$ as the transition function, and $F \subseteq Z$ as the set of final states.

The NFA for q_γ is obtained by the *Thompson construction* [55] that recursively builds small automaton fragments and links them with ε -transitions. For each sub-expression (concatenation, union, Kleene star), a fragment is created with unique initial and accept states: union merges sub-expressions via a new initial and accept state, concatenation joins accept states of one fragment to the start state of the next, and Kleene star introduces a loop to re-enter the sub-expression or skip it. By recursively assembling these fragments, one obtains an NFA accepting all words in $L(q_\gamma)$.

As motivated, we aim to detect *all possible (partial) matches of words from $L(q_\gamma)$ to indices in a word w* . This mirrors the most challenging combination of consumption and selection policies [3, 25, 58] in event processing: *Reuse* and *Skip-Till-Any-Match (STAM)*. A simple automata-based matching method works as follows: Each incoming element is checked against the current set of automata. If an element enables a transition in an automaton, the automaton is duplicated, with one copy disregarding the transition and the other one taking it, yielding an exponential growth in the number of automata as the input grows.

Regular Expression Subsequence Matching. To formalize the evaluation of a regular expression query $q = (\gamma, \tau, \mathcal{A})$, only q_γ and q_τ are required, as the aggregate function $\mathcal{A}$ is applied to the resulting matches. We capture the matches resulting from evaluating q over a stream $s = (\omega, \lambda, \nu)$ using two sets: complete matches CM and partial matches PM . These sets are based on the language $L(q_\gamma)$ and its prefix language $L_{\text{pre}}(q_\gamma)$, respectively. PM includes a partial match if it maps a word $y \in L_{\text{pre}}(q_\gamma)$ to indices in ω . CM includes a match if it maps a word $y \in L(q_\gamma)$ to indices in ω . Either way, the mapping must satisfy the time window q_τ .

Next, we formalize the construction of CM , while the construction of PM works analogously. Let i be the index of the last received stream element e_i and $[i] = \{1, \dots, i\}$. We define $\mathcal{M}_{\text{CM}} = \{\varphi \mid \varphi: L(q_\gamma) \rightarrow 2^{[i]}\}$ as the set of all *word mapping functions* φ from words in $L(q_\gamma)$ to sets of up to i indices in ω . Note that the mapping function m introduced earlier can be depicted through the word mapping function φ_m , e.g., in Example 2, m_1 aligns with $\varphi_{m_1}(\langle A_1, B_2, C_3 \rangle) = \{1, 2, 3\}$. Analogously, the set of all *partial word mapping functions*, $\mathcal{M}_{\text{PM}} = \{\varphi' \mid \varphi': L_{\text{pre}}(q_\gamma) \setminus L(q_\gamma) \rightarrow 2^{[i]}\}$, is based on $L_{\text{pre}}(q_\gamma)$.

Table 1. Overview of notations for RegEx queries.

Notation	Explanation
$\Sigma = \{a_1, \dots, a_m\}$	Alphabet containing characters.
$w = \langle c_1, \dots, c_n \rangle$	Word (sequence of characters) of size $n = w $.
$w[i]$	Character at i -th index of w .
$y \leq_m w$	A subsequence y of w with a mapping m .
$q = (\gamma, \tau, \mathcal{A})$	Regular expression query q , composed of regular expression γ , a time window size τ , and an aggregate function $\mathcal{A}$.
$s = (\omega = \langle c_1, \dots \rangle,$ $\lambda = \langle t_1, \dots \rangle,$ $v = \langle v_1, \dots \rangle)$	A <i>stream</i> of related sequences of characters $c_i \in \Sigma$ (i.e., ω), timestamps $t_i \in \mathbb{T}$ (i.e., λ), and values $v_i \in \mathbb{N}$ (i.e., v), respectively.
$L(\gamma), L_{\text{pre}}(\gamma)$	The languages of <i>all words</i> and <i>all prefixes</i> of γ , respectively.
PM, CM	The sets of partial and complete matches, respectively, resulting from evaluating q over the stream s .

A word mapping function φ for a word $y \in L(q_\gamma)$, represented by $\varphi(y) = \{i_1, i_2, \dots, i_{|y|}\}$ with $i_1 < i_2 < \dots < i_{|y|}$, is deemed *valid* if it satisfies the following conditions:

- (1) (Reconstruction) The indices $\varphi(y) = \{i_1, i_2, \dots, i_{|y|}\}$ in ω must reconstruct $y \in L(q_\gamma)$, i.e., $y = \langle \omega[i_1], \omega[i_2], \dots, \omega[i_{|y|}] \rangle$.
- (2) (Time Window Constraint) The mapping must adhere the time window constraint, i.e., $\lambda[i_{|y|}] - \lambda[i_1] \leq q_\tau$.

Based thereon, we define the set $CM = \{\varphi \mid \varphi \in \mathcal{M}_{CM}, \varphi \text{ is valid}\}$ as all valid word mapping functions φ . Analogously, PM denotes the set of all valid partial word mapping functions φ' .

Finally, using the set of complete matches as CM , we integrate the aggregate function $\mathcal{A}$. With $q_{\mathcal{A}} = \text{COUNT}$, the result is $|CM|$, i.e., the total number of complete matches. If $q_{\mathcal{A}} = \text{SUM}$, the result is $\sum_{\varphi \in CM} \sum_{j \in \text{range}(\varphi)} v[j]$, summing values from v of complete matches. For $q_{\mathcal{A}} = \text{AVG}$, the result is $\text{SUM}(CM)/\text{COUNT}(CM)$.

EXAMPLE 3. Consider an alphabet $\Sigma = \{A, B, C\}$ and a regular expression query $q = (\gamma, \tau, \mathcal{A})$ with $q_\gamma = ABC$ and $q_\tau = 5$. The language of q_γ is $L(q_\gamma) = \{\langle A_1, B_2, C_3 \rangle\}$. For the stream $s = (\omega, \lambda, v)$, with $\omega = \langle A_1, C_2, B_3, A_4, C_5, B_6, C_7 \rangle$, $\lambda = \langle 1, 2, 3, 4, 5, 6, 7 \rangle$, and values $v = \langle 2, 3, 5, 6, 13, 7, 9 \rangle$, potential matches are $\{\varphi_1: \langle A_1, B_2, C_3 \rangle \mapsto \{1, 3, 5\}, \varphi_2: \langle A_1, B_2, C_3 \rangle \mapsto \{1, 3, 7\}, \varphi_3: \langle A_1, B_2, C_3 \rangle \mapsto \{4, 6, 7\}\}$. Incorporating the time window q_τ , we obtain $CM = \{\varphi_1: \langle A_1, B_2, C_3 \rangle \mapsto \{1, 3, 5\}, \varphi_3: \langle A_1, B_2, C_3 \rangle \mapsto \{4, 6, 7\}\}$. For some aggregations $q_{\mathcal{A}}$, we get $\text{COUNT}(CM) = 2$ and $\text{SUM}(CM) = (2+5+13) + (6+7+9) = 42$.

4 Problem Statement

We consider scenarios, in which the query evaluation is relevant at a random time instance (e.g., when a user refreshes a dashboard). Let $U: \Omega \rightarrow \mathbb{T}$ be a random variable representing this trigger, with Ω as the sample space encompassing all potential trigger instances. The probability distribution of U is denoted by $P_U(t)$, specifying the probability that U equals t . Let $\mathcal{E} \subseteq \mathbb{T}$ be the set of *evaluation timestamps*, which are drawn independently as a fixed number N of samples from $\mathbb{T}$ according to P_U , yielding $\mathcal{E} = \{\varepsilon_1, \dots, \varepsilon_N\}$. Each $\varepsilon_i \in \mathcal{E}$ denotes a time instance when query results are relevant.

Then, we define a summary substream $\mathcal{S}$ as a subsequence of elements from the substream $s(\varepsilon)$ up to time $\varepsilon \in \mathcal{E}$, containing at most n elements, as $\mathcal{S}(\varepsilon, n) = (\langle c_{i_1}, \dots, c_{i_k} \rangle, \langle t_{i_1}, \dots, t_{i_k} \rangle, \langle v_{i_1}, \dots, v_{i_k} \rangle)$ where $1 \leq i_1 \leq \dots \leq i_k \leq p$, $k \leq n$, and each i_j is an index of a stream element in $s(\varepsilon)$. Here, p is the index of the latest stream element with timestamp $\leq \varepsilon$. Technically, the summary substream

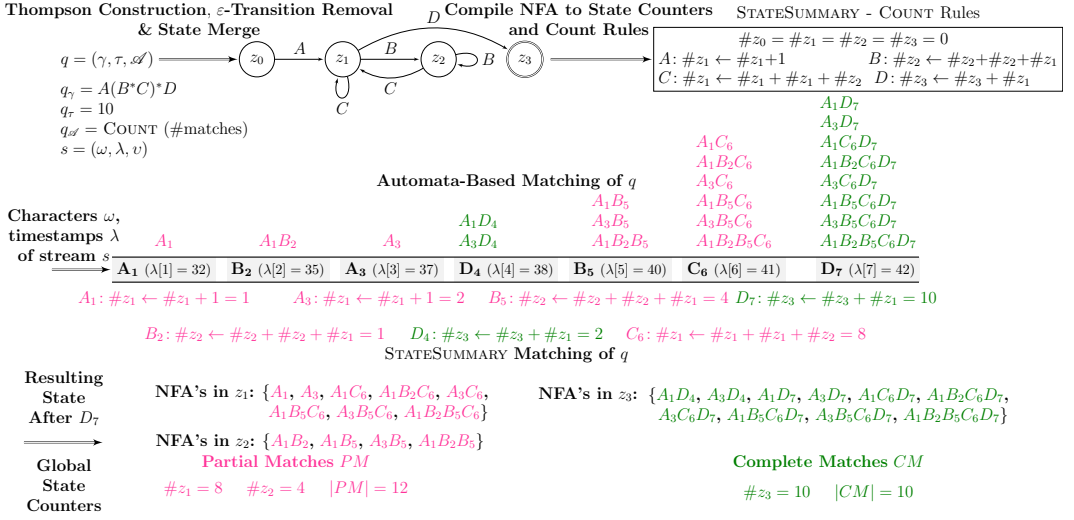


Fig. 3. Characters ω and timestamps λ of a stream s and the NFA for the regular expression $q_\gamma = A(B^*C)^*D$. The counters of the STATESUMMARY are contrasted with the (partial) matches constructed by traditional automata-based matching.

is constructed by a summary selection function that decides whether an element should be kept, discarded, or replaced.

We aim to assess how the projected elements in a summary affect a loss function $l_{\mathcal{A}} : \mathbb{N} \times \mathbb{N} \rightarrow \mathbb{R}$ upon an evaluation trigger. The loss quantifies the discrepancy between the values x_{opt} and x , obtained by the aggregate for query matches over either the stream $s(\varepsilon)$ (for x_{opt}) or the chosen subsequence $\mathcal{S}(\varepsilon, n)$ (for x) at time ε .

For aggregations COUNT and SUM, we employ the loss function $l_{\text{COUNT}}(x, x_{\text{opt}}) = l_{\text{SUM}}(x, x_{\text{opt}}) = |x - x_{\text{opt}}|$. For AVG, we define the loss as the relative error $l_{\text{AVG}}(x, x_{\text{opt}}) = \left| \frac{x - x_{\text{opt}}}{x_{\text{opt}}} \right|$. For each evaluation timestamp $\varepsilon \in \mathcal{E}$, we aim to minimize the loss $l_{\mathcal{A}}$.

PROBLEM 1. Let Σ be an alphabet and $s(\varepsilon) = (\omega, \lambda, v)$ be a stream up to time instance $\varepsilon \in \mathbb{T}$. Also, let $q = (\gamma, \tau, \mathcal{A})$ be a regular expression query, n be the summary size, $\mathcal{S}(\varepsilon, n)$ the summary substream, and $\mathcal{E} \subseteq \mathbb{T}$ a finite set of evaluation timestamps. For each time point $\varepsilon \in \mathcal{E}$, we aim at a summary substream $\mathcal{S}(\varepsilon, n)$ with minimal loss:

$$\text{minimize} \sum_{\varepsilon \in \mathcal{E}} \min_{\mathcal{S}(\varepsilon, n)} (l_{\mathcal{A}}(q_{\mathcal{A}}(CM_{\mathcal{S}(\varepsilon, n)}), q_{\mathcal{A}}(CM_{s(\varepsilon)})))$$

5 State Summary Data Structure

The optimization problem defined in Problem 1 seeks to minimize the loss at evaluation time points, which, in turn, is affected by the set of complete matches at these time points. However, as the time points are unknown, the complete matches would need to be available at *any* time. Yet, their computation at every time point is infeasible. As mentioned, the number of automata to maintain may grow exponentially in the number of processed elements, which yields an exponential runtime of an evaluation algorithm.

Motivated by these observations, we present SuSe to address Problem 1. We avoid any direct materialization and maintenance of (partial) matches and retain only the information necessary for evaluating the aggregate function, at any time point. Specifically, SuSe relies on a stream

summary, as illustrated in Fig. 1. The summary captures aggregated stream information that, when an evaluation trigger occurs, allows for the (approximate) computation of the aggregate over the complete matches of the query.

The stream summary employed in SuSE comprises the summary substream of selected stream elements $\mathcal{S}$ and the STATESUMMARY data structure. Given a regular expression query $q = (\gamma, \tau, \mathcal{A})$, we compile the STATESUMMARY from it, which stores aggregated information about partial and complete matches from evaluating q over elements in $\mathcal{S}$. We refer to these aggregates as *global state counters*: For each state of the underlying NFA of q_γ , a global state counter tracks the aggregated information for (partial) matches in that state, considering all elements that have been selected into the summary substream. To measure how each element in $\mathcal{S}$ contributes to these global state counters, we introduce *local state counters*, capturing for each state the aggregated information that particular element contributes to. In §6, local state counters guide our summary selection algorithm in deciding which elements to keep or replace by assigning them a *benefit*. Finally, to handle a temporal context due to the time window, we define an *active time window* that denotes the segment of summary elements still capable of producing future matches. We use this segment to enrich the potential benefit of summary elements by predicting their future match participation.

In the remainder, we fix a time instance $\varepsilon \in \mathbb{T}$ and summary size n , and, thus, write $\mathcal{S}$ to represent the summary substream. Let $s = (\omega, \lambda, v)$ be a stream; the summary substream $\mathcal{S} = (\omega', \lambda', v')$ is a subsequence of its elements.

To simplify the presentation, we first introduce the STATESUMMARY for queries using COUNT as the aggregation function, in terms of its construction (§5.1) and maintenance (§5.2). Then, we discuss how it is adapted for SUM and AVG aggregations (§5.3).

5.1 STATESUMMARY Compilation

Given a regular expression query $q = (\gamma, \tau, \mathcal{A})$, the STATESUMMARY is constructed based on the NFA for q_γ obtained by the Thompson construction [55]. The NFA is further adapted by removing epsilon transitions, subsequent removal of unreachable states, and merging of states (which are final or non-final) with identical transitions. The resulting NFA, denoted as $\mathcal{N} = (Z, \Sigma, \delta, z_0, F)$, is used to compile our STATESUMMARY. The data structure includes state counters (summarized in Table 2) that are defined and derived as follows.

Global State Counter. For each state $z_i \in Z$ in $\mathcal{N}$, where $0 \leq i < |Z|$, the *global state counter* $\#z_i \in \mathbb{N}$ keeps track of the number of (partial) matches in that state z_i , considering *all* selected elements $e_i = (\omega'[i], \lambda'[i], v'[i])$ in $\mathcal{S}$. Initially, these counters are set to zero. Let $\mathcal{S}_{\#Z} = (\#z_0, \dots, \#z_{|Z|})$ denote the tuple of *global state counters* of $\mathcal{S}$, derived from the states Z of the compiled NFA, where each component $\#z_i$ in $\mathcal{S}_{\#Z}$ denotes a global state counter. Referring to the example given in Fig. 3, with global state counters $(\#z_0, \#z_1, \#z_2, \#z_3)$, processing the stream s yields 8 partial matches in z_1 , 4 partial matches in z_2 , and 10 complete matches in z_3 , represented by the global state counters $(0, 8, 4, 10)$.

COUNT Rules. The COUNT rules specify the number of (partial) matches in a given state, based on the evaluation of query q over the selected elements in $\mathcal{S}$. We derive them from the transitions of $\mathcal{N}$ by checking, for each state $z_i \in Z$ and each character $c \in \Sigma$, if there is a transition leading to any state in Z . As an illustration, consider the NFA in Fig. 3, where $\delta(z_1, C) = \{z_1\}$, while $\delta(z_1, A) = \emptyset$.

With each character c , any (partial) match in a state z_i satisfying $\delta(z_i, c) \neq \emptyset$ can transition using c . This creates (partial) matches in the resulting state(s), as each (partial) match in z_i transitions with c , causing the current count of (partial) matches in the subsequent state(s), given by $\delta(z_i, c)$, to increment by the current count $\#z_i$. Transitions from the initial state $z_0 \in Z$ to a state $z_k \neq z_0$ represent the creation of a new (partial) match, adding one to $\#z_k$.

Table 2. Overview of different counter concepts.

Notation	Explanation
$\mathcal{S}_{\#Z}$	Global state counters: counting <i>all</i> (partial) matches in $\mathcal{S}$.
$e_{\#Z}$	Local state counters: counting (partial) matches e is involved.
$\mathcal{A}_{\#Z}$	Active global state counters: counting <i>all</i> (partial) matches satisfying q_τ , thus having the potential to lead to future matches.
$\mathcal{A}_{e_{\#Z}}$	Active local state counters: counting (partial) matches e is involved satisfying q_τ , thus, having the potential to lead to future matches.

Based on $\mathcal{N}$'s transitions, we define COUNT rules to quantify the number of (partial) matches within a state resulting from matching an element against the current state. For each state $z_i \in Z$, global state counter $\#z_i$ in $\mathcal{S}_{\#Z}$, and character $c \in \Sigma$:

- (1) If z_i has a self-loop with c :
if $\delta(z_i, c) = \{z_i\}$ then $\#z_i \leftarrow \#z_i + \#z_i$.
- (2) If transitions from states $\{z_{i_1}, z_{i_2}, \dots, z_{i_l}\}$ lead to z_k upon processing character c (excluding the self-loop case):
if $\delta(z_{i_j}, c) = \{z_k\}$ for $j \in \{1, \dots, l\}$ then $\#z_k \leftarrow \#z_k + \sum_{j=1}^l \#z_{i_j}$.
- (3) If the transition goes from the initial state z_0 to $z_k \neq z_0$:
if $\delta(z_0, c) = \{z_k\}$ then $\#z_k \leftarrow \#z_k + 1$.

Note that for counters impacted by both, a self-loop and a transition from a different state, the number must first be doubled by the self-loop. This ensures that no (partial) match undergoes two transitions.

Since δ yields a subset of 2^Z , a character c may update multiple state counters, with $O(|Z|)$ as an upper bound.

EXAMPLE 4. Consider the stream $s = (\omega, \lambda, v)$ with $\omega = \langle A_1, B_2, A_3, D_4, B_5, C_6, D_7 \rangle$ in Fig. 3. The expression $\gamma = A(B^*C)^*D$ yields global state counters $\mathcal{S}_{\#Z} = (\#z_0, \#z_1, \#z_2, \#z_3)$, each initialized to zero. The COUNT rules are depicted in the upper box in Fig. 3. Processing s , A_1 increments $\#z_1$ by one. Then, B_2 increases $\#z_2$ by the sum of $\#z_1$ and $\#z_2$; A_3 adds one to $\#z_1$; and D_4 augments $\#z_3$ to 2. After processing s , it holds that $|PM| = \#z_1 + \#z_2 = 8 + 4 = 12$ and $|CM| = \#z_3 = 10$.

Local State Counter. For each element e in $\mathcal{S}$, *local state counters*, denoted as $e_{\#Z} = (e_{\#z_0}, \dots, e_{\#z_{|Z|}})$ and initialized to zero, quantify the number of partial and complete matches involving element e across each state. They are updated using the COUNT rules introduced above for global state counters. When an element e arrives, the local state counters $e_{\#Z}$ reflect the (partial) matches that e contributes to the global state counters; thus, $e_{\#Z}$ is updated first based on the values of the global state counters in $\mathcal{S}_{\#Z}$. Subsequent updates of $e_{\#Z}$, however, rely on the values in their respective local state counters, to incorporate new (partial) matches involving e .

EXAMPLE 5. Consider Fig. 3 and element B_5 . Upon its arrival, B_5 is incorporated by all partial matches in states z_1 and z_2 . Hence, we have $B_{5\#z_2} \leftarrow \#z_2 + \#z_1 = 3$. Note that only when B_5 is added, we apply the global state counters $\#z_i$ in $\mathcal{S}_{\#Z}$ to the local state counters. Afterwards, the COUNT rules update B_5 's local state counters by factoring in the corresponding local state counter values. Therefore, when C_6 is processed next, the local state counter for B_5 gets updated: $B_{5\#z_1} \leftarrow B_{5\#z_1} + B_{5\#z_1} + B_{5\#z_2} = 3$, resulting in $B_{5\#Z} = (0, 3, 3, 0)$.

To discuss complexity, we overload notation and write $|\mathcal{S}|$ for the number of elements in $\mathcal{S}$. If each state is reachable from any other state, for each element e , we must monitor all $|Z|$ counters,

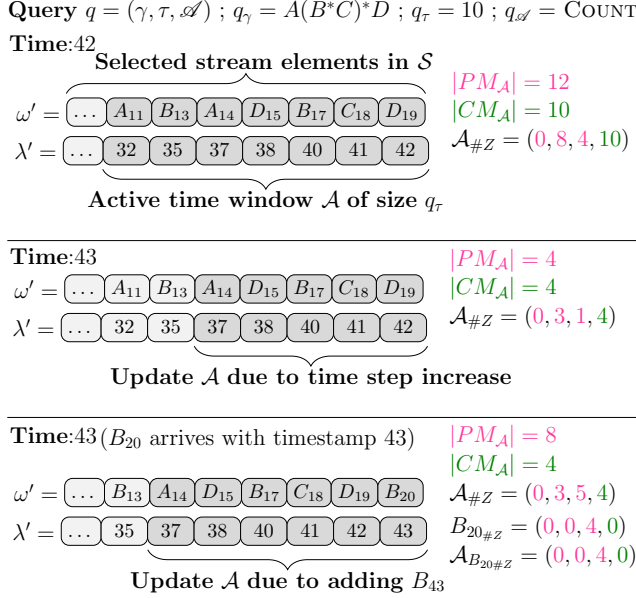


Fig. 4. Evolution of a summary for a given regular expression query; shaded elements are in the active time window $\mathcal{A}$.

which yields a space complexity for both types of counters of $O(|Z| + |S| \cdot |Z|)$, which simplifies to $O(|S| \cdot |Z|)$, where typically $|Z| \ll |S|$.

Sliding Time Window. The time window q_τ of a regular expression query constrains matches based on the timestamps of elements. If a stream element e is added to the summary, it is possible that not all summary elements fall within the same time window as e . Hence, the global state counters $\mathcal{S}_{\#Z}$ may not represent the correct values for initializing the local state counters $e_{\#Z}$. We now detail the extensions to STATESUMMARY for such temporal context.

We define the *active time window* $\mathcal{A}$ as the subsequence of the selected elements in $\mathcal{S}$ that, upon adding an element e , satisfy the time window q_τ with e . Based thereon, we maintain separate *active* global and local state counters, denoted by $\mathcal{A}_{\#Z}$ and $\mathcal{A}_{e\#Z}$, respectively. They reflect the match counts per element in $\mathcal{A}$ and are updated based on the same COUNT rules as before. Within $\mathcal{A}$, the term $\mathcal{A}_{init}$ denotes the oldest initiator element, i.e., the oldest selected element, for which the character aligns with a transition in the automaton starting in the initial state z_0 . There is no need to keep elements before this element $\mathcal{A}_{init}$ not linked to transitions from the initial state, as they cannot create *new* (partial) matches.

EXAMPLE 6. Consider the selected characters ω' of the elements at time 42 in Fig. 4 and the query q from Fig. 3. Here, D_{19} is the most recently appended element. Due to q_τ , $A_{11} = \mathcal{A}_{init}$ remains the earliest initiator matching D_{19} and stays in $\mathcal{A}$ as it satisfies the time window. Yet, at time 43, A_{11} is no longer valid and, hence, is removed from $\mathcal{A}$, but it remains in $\mathcal{S}$. Although B_{13} fits within the time window, it is discarded from $\mathcal{A}$ since it is not an initiator; we set $\mathcal{A}_{init} = A_{14}$. The elements in $\mathcal{A}$ correspond to ω in Fig. 2; thus, the matches corresponding to the counters are also depicted there.

As we add B_{20} , the local state counters $B_{20\#Z}$ are initialized from the active time window state counters $\mathcal{A}_{\#Z} = (0, 3, 1, 4)$, so that we arrive at $B_{20\#Z} = (0, 0, \mathcal{A}_{\#z_1} + \mathcal{A}_{\#z_2}, 0) = (0, 0, 4, 0)$. These counts represent the current number of (partial) matches. Then, we update the global state counter values by adding the contributions from B_{20} : $\mathcal{S}_{\#Z} \leftarrow \mathcal{S}_{\#Z} + B_{20\#Z}$, meaning $\mathcal{S}_{\#z_i} \leftarrow \mathcal{S}_{\#z_i} + B_{20\#z_i}$ for $0 \leq i < |Z|$.

In addition, we maintain *active local state counters* $\mathcal{A}_{e_{\#Z}}$ for elements e in $\mathcal{A}$ to track the number of matches each element participates in within $\mathcal{A}$ (whereas the local state counters $e_{\#Z}$ also cover elements no longer in $\mathcal{A}$). This nuanced tracking is crucial for estimating the benefit of an element in $\mathcal{A}$ to yield future matches.

As soon as an element e drops out of $\mathcal{A}$, its local state counters $e_{\#Z}$ can no longer increase, but only decrease. The reason being that e can no longer lead to new (partial) matches.

Maintaining the active global and local state counters results in additional space requirements of size $O(|\mathcal{A}| \cdot |Z|)$. Since $|\mathcal{A}|$ depends on q_τ and $|\mathcal{S}|$, we arrive at $O(\min\{q_\tau, |\mathcal{S}|\} \cdot |Z|)$.

5.2 Operations

In this section, we detail the INITIALIZE, INSERT, and REMOVE operations that manipulate the summary.

INITIALIZE. The STATESUMMARY is initialized for a regular expression query $q = (\gamma, \tau, \mathcal{A})$, the NFA $\mathcal{N} = (Z, \Sigma, \delta, z_0, E)$ derived for it, and a user-defined *summary size* n (from Problem 1). We create tuples $\mathcal{S}_{\#Z}, \mathcal{A}_{\#Z}$ to capture the global state counters and the active global state counters, respectively. Note that the active time window $\mathcal{A}$ is a subsequence of $\mathcal{S}$, so that we do not store it explicitly.

We allocate $O(n \cdot |Z|)$ memory for all future elements and their local state counters. Consequently, the runtime of the INITIALIZE operation is dominated by initializing the local state counter tuples for these elements, which requires $O(|\mathcal{S}| \cdot |Z|)$ time and space.

We use an array to hold the elements when implementing the STATESUMMARY. While this choice induces an overhead for the REMOVE operation, it facilitates efficient linear scans, which are required by our INSERT and REMOVE operations.

INSERT. The INSERT(e) operation appends an element e from the stream to the summary substream $\mathcal{S}$ and is detailed in Alg. 1. To incorporate e , we need to adjust the (active) global state counters $\mathcal{S}_{\#Z}$ and $\mathcal{A}_{\#Z}$, and for each element $\hat{e}$ in $\mathcal{A}$ its (active) local state counters $\hat{e}_{\#Z}$ and $\mathcal{A}_{\hat{e}_{\#Z}}$. Finally, we have to initialize the active local state counter $\mathcal{A}_{e_{\#Z}}$ of e and its local state counters $e_{\#Z}$.

To comprehend the influence of e 's character c on the NFA, we assess its impact on ongoing (partial) matches within the active time window $\mathcal{A}$. Specifically, we identify the states affected by c . This is determined by the COMPUTESTATECOUNTERCHANGE function (line 11), which first initializes a new state counter tuple *newCounters* of size $|Z|$ with zeros. Then, for each source state *from_z* $\in Z$, we determine each target state *to_z* $\in Z$, resulting from a transition using c , to increase the state counter *newCounters*_{#to_z} by the count of *counters*_{#from_z} (line 15). The resulting update tuple is saved to *globCounterChange* (line 1).

In line 2 and line 3, we update the global state counters $\mathcal{S}_{\#Z}$ and active global state counters $\mathcal{A}_{\#Z}$ using *globCounterChange* (i.e., by adding tuple components, line 17). From line 4 to line 7, for each element $\hat{e}$ in $\mathcal{A}$, we determine the impact of c on the active local state counters of $\hat{e}$, denoted by *localChange* (line 5), to update its active local state counters $\mathcal{A}_{\hat{e}_{\#Z}}$ and local state counters $\hat{e}_{\#Z}$ by *localChange*. Finally, after appending the new element e to $\mathcal{S}$ (line 8), denoted by $\mathcal{S} \oplus e$, we initialize its active local state counters $\mathcal{A}_{e_{\#Z}}$ (line 9) and local state counters $e_{\#Z}$ (line 10).

The time complexity of INSERT is dominated by the updates of the (active) local state counters of elements in $\mathcal{A}$. The size of $\mathcal{A}$ is dictated by q_τ and, if $q_\tau > |\mathcal{S}|$, a single insert can affect *all* elements in $\mathcal{S}$. Since COMPUTESTATECOUNTERCHANGE requires $O(|Z|^2)$ time, INSERT takes $O(\min\{q_\tau, |\mathcal{S}|\} \cdot |Z|^2)$ time using $O(|Z|)$ space.

REMOVE. The REMOVE(e) operation removes element e from $\mathcal{S}$. From global state counters $\mathcal{S}_{\#Z}$, we subtract the local state counters $e_{\#Z}$, representing all (partial) matches involving e . Similarly, the active global state counters $\mathcal{A}_{\#Z}$ are reduced by $\mathcal{A}_{e_{\#Z}}$. When removing e , related (active) local

Algorithm 1: INSERT.

input :Summary substream S ; global state counters $S_{\#Z}$ and all local state counters (e.g., $e_{\#Z}$); stream element e to insert and its character c ; NFA $N = (Z, \Sigma, \delta, z_0, E)$; active time window $\mathcal{A}$ with active global state counters $\mathcal{A}_{\#Z}$ and all active local state counters (e.g., $\mathcal{A}_{e_{\#Z}}$)

```

1 globCounterChange  $\leftarrow$  computeStateCounterChange( $\mathcal{A}_{\#Z}, N, c$ );
2  $S_{\#Z} \leftarrow$  addStateCounters( $S_{\#Z}, \text{globCounterChange}, N$ );
3  $\mathcal{A}_{\#Z} \leftarrow$  addStateCounters( $\mathcal{A}_{\#Z}, \text{globCounterChange}, N$ );
4 for  $\hat{e}$  in  $\mathcal{A}$  do
5   localChange  $\leftarrow$  computeStateCounterChange( $\mathcal{A}_{\hat{e}_{\#Z}}, N, c$ );
6    $\mathcal{A}_{\hat{e}_{\#Z}} \leftarrow$  addStateCounters( $\mathcal{A}_{\hat{e}_{\#Z}}, \text{localChange}, N$ );
7    $\hat{e}_{\#Z} \leftarrow$  addStateCounters( $\hat{e}_{\#Z}, \text{localChange}, N$ );
8  $S \leftarrow S \oplus e$ ;
9  $\mathcal{A}_{e_{\#Z}} \leftarrow$  globCounterChange;
10  $e_{\#Z} \leftarrow$  globCounterChange;
11 function computeStateCounterChange( $counters, N, c$ )
12   newCounters  $\leftarrow 0^{|Z|}$ ;
13   for  $from\_z \in Z$  do
14     for  $to\_z \in \delta(from\_z, c)$  do
15       newCounters $_{\#to\_z} \leftarrow$  newCounters $_{\#to\_z} +$  counters $_{\#from\_z}$ ;
16   return newCounters;
17 function addStateCounters( $stateCounters1, stateCounters2, N$ )
18   addCounters  $\leftarrow 0^{|Z|}$ ;
19   for  $z \in Z$  do addCounters $_{\#z} \leftarrow$  stateCounters1 $_{\#z} +$  stateCounters2 $_{\#z}$ ;
20   return addCounters;

```

Table 3. Time and space complexity of summary operations.

Operation	Time complexity	Space complexity
INITIALIZE	$O(S \cdot Z)$	$O(S \cdot Z)$
INSERT	$O(\min\{q_\tau, S \} \cdot Z ^2)$	$O(Z)$
REMOVE	$O(\min\{q_\tau, S \}^2 \cdot Z ^2)$	$O(\min\{q_\tau, S \} \cdot Z)$

state counters for remaining elements may require updating. However, STATESUMMARY does not include the required counts of (partial) matches for pairs of elements. We, thus, define a procedure for obtaining the joint counts of e and all affected elements, as detailed in our technical report [6].

Complexity. REMOVE implies a series of INSERT operations. For every new element inserted, all prior elements are updated. This results in a quadratic time complexity of $O(\min\{q_\tau, |S|\}^2 \cdot |Z|^2)$ and a space complexity of $O(\min\{q_\tau, |S|\} \cdot |Z|)$. We give an overview of our operations in Table 3 and refer to [6] for details.

5.3 Model Extensions

Aggregate Functions. Having described the STATESUMMARY for COUNT aggregations, we turn to SUM and AVG.

As before, when an element e is added to S , transitions of the automaton are triggered, creating new matches in subsequent states. For all affected (partial) matches, the character c of element e and its payload value v must be incorporated. We define the SUM counters and rules based on those introduced for COUNT.

For each state $z_i \in Z$, we create a SUM counter $S_{\#z_i} \in \mathbb{N}$, representing the current total sum over all payload values of (partial) matches in z_i . This yields a tuple of *global sum counters* $S_{\#Z} =$

$(\mathcal{S}_{\#z_0}, \dots, \mathcal{S}_{\#z_{|Z|}})$. As for the COUNT rules, there are three cases we have to handle for defining the SUM rules: (1) self-loops, (2) transitions from other states, and (3) transitions originating from z_0 : For each $z_i \in Z$, each $\mathcal{S}_{\#z_i}$ in $\mathcal{S}_{\#Z}$, and character c with value v :

- (1) If z_i has a self-loop with c :
if $\delta(z_i, c) = \{z_i\}$ then $\mathcal{S}_{\#z_i} \leftarrow \mathcal{S}_{\#z_i} + \#z_i \cdot v$.
- (2) If transitions from states $\{z_{i_1}, z_{i_2}, \dots, z_{i_l}\}$ lead to z_k upon processing character c (excluding the self-loop case):
if $\delta(z_{i_j}, c) = \{z_k\}$ for $j \in \{1, 2, \dots, l\}$ then
$$\mathcal{S}_{\#z_k} \leftarrow \mathcal{S}_{\#z_k} + \sum_{j=1}^l \mathcal{S}_{\#z_{i_j}} + \#z_{i_j} \cdot v$$
- (3) If the transition goes from the initial state z_0 to $z_k \neq z_0$:
if $\delta(z_0, c) = \{z_k\}$ then $\mathcal{S}_{\#z_k} \leftarrow \mathcal{S}_{\#z_k} + 1 \cdot v$.

The (active) local sum counters are defined analogously to their COUNT-based counterparts, but rely on the SUM rules.

Maintaining the STATESUMMARY for SUM and COUNT, support for AVG is obtained by using the ratio $\mathcal{S}_{\#z_i} / \#z_i$ for each state $z_i \in Z$.

EXAMPLE 7. *Extending our example from §2, with a stream partitioned by bike ID and a query detecting all short bike trips, each trip event may include a duration attribute. Evaluating the SUM aggregate on this attribute yields the total ride time per bike across all such trips, whereas AVG computes the avg ride time per bike per trip.*

Predicates. Patterns often materialize in stream partitions (e.g., per sensor, per person, per vehicle), making our model directly applicable for equality predicates over corresponding attribute values per partition. When defining predicates over subsets of query characters, e.g., $q_r = ABC$ where B 's and C 's (but not A 's) shall have equal attribute values, characters for which the predicate does not apply (here, the A characters) are replicated in each partition. Without partitioning, equality predicates can still be processed by maintaining counters for each attribute value previously considered for partitioning, tracking matches specifically for that value.

Selection Policies. The semantics underlying our model is known as the STAM policy [3, 25, 58]. To integrate a stricter policy, such as Skip-Till-Next-Match (STNM) [25], we adapt the COUNT rules to reflect STNM's key distinction: it prevents skipping *relevant* elements, i.e., (partial) matches are not duplicated; thus, the increase in matches in subsequent states is balanced by an equivalent decrease in the originating states when an element arrives.

Multi-Query. When multiple RegEx queries share the same time window size τ , they can be processed together with one STATESUMMARY by merging their patterns via union, e.g., $q_1 = (\gamma_1, \tau, \mathcal{A})$ and $q_2 = (\gamma_2, \tau, \mathcal{A})$ become $q_3 = (\gamma_1 | \gamma_2, \tau, \mathcal{A})$. If queries have different time window sizes, each STATESUMMARY must manage distinct active time window segments, complicating sharing and maintenance. Still, a separate STATESUMMARY for each query may be used, though it forfeits shared-counter benefits and increases maintenance overhead by at least a factor in the size of the query workload.

6 Summary Selection

To maintain a summary with the aim to minimize the loss in query evaluation, we first show how to quantify the potential of an element to contribute to a substantial amount of complete matches (§6.1). Based thereon, we present a selection strategy (§6.2). Again, we first discuss COUNT, and later generalize to SUM and AVG (§6.3).

Query:	Alph. prob.:	$t = 0$	$t = 1$
$q = (\gamma, \tau, \mathcal{A})$	$P(A) = 0.5,$	$\#z_1(0) = \#z_1$	$\#z_1(1) = \#z_1(0) + \overbrace{1 \cdot P(A) + (\#z_1(0) + \#z_2(0)) \cdot P(C)}^{\text{Weighted influences of } A \text{ and } C \text{ at } t=0 \text{ on } \#z_1}$
$q_\gamma = A(B^*C)^*D$	$P(B) = 0.2,$	$\#z_2(0) = \#z_2$	$\#z_1(1) = \#z_1 + 0.5 + \#z_1 \cdot 0.15 + \#z_2 \cdot 0.15$
$q_\tau = \tau$	$P(C) = 0.15,$	$\#z_3(0) = \#z_3$	$\#z_1(1) = 1.15 \cdot \#z_1 + 0.15 \cdot \#z_2 + 0.5$
$q_{\mathcal{A}} = \text{COUNT}$	$P(D) = 0.15$		$I_1(0) = 1 \cdot P(A)$
$\#z_0 = \#z_1 = \#z_2 = \#z_3 = 0$			$I_4(0) = \#z_1(0) \cdot P(C)$
$A: \#z_1 \leftarrow \#z_1 + 1$			$I_5(0) = \#z_2(0) \cdot P(C)$
$B: \#z_2 \leftarrow \#z_2 + \#z_2 + \#z_1$			$\#z_2(1) = \#z_2(0) + (\#z_2(0) + \#z_1(0)) \cdot P(B)$
$C: \#z_1 \leftarrow \#z_1 + \#z_1 + \#z_2$			$\#z_2(1) = \#z_2 + \#z_2 \cdot 0.2 + \#z_1 \cdot 0.2$
$D: \#z_3 \leftarrow \#z_3 + \#z_1$			$I_2(0) = \#z_2(0) \cdot P(B)$
Influences on $\#z_1$: $1 + \#z_1 + \#z_2$			$I_3(0) = \#z_1(0) \cdot P(B)$
Influences on $\#z_2$: $\#z_2 + \#z_1$			$\#z_2(1) = 1.2 \cdot \#z_2 + 0.2 \cdot \#z_1$
Influence on $\#z_3$: $\#z_1$			$\#z_3(1) = \#z_3(0) + \#z_1(0) \cdot P(D)$
			$\#z_3(1) = \#z_3 + 0.15 \cdot \#z_1$
			$I_6(0) = \#z_1(0) \cdot P(D)$
		$t = 2$	$t = 1$
			Weighted influences of A and C at $t=1$ on $\#z_1$
			$\#z_1(2) = \#z_1(1) + \overbrace{1 \cdot P(A) + (\#z_1(1) + \#z_2(1)) \cdot P(C)}^{\text{Weighted influences of } A \text{ and } C \text{ at } t=1 \text{ on } \#z_1}$
			$\#z_1(2) = 1.15 \cdot \#z_1 + 0.15 \cdot \#z_2 + 0.5 + P(A) +$
			$(1.15 \cdot \#z_1 + 0.15 \cdot \#z_2 + 0.5 + 1.2 \cdot \#z_2 + 0.2 \cdot \#z_1) \cdot P(C)$
			$\#z_1(2) = 1.15 \cdot (1.15 \cdot \#z_1 + 0.15 \cdot \#z_2 + 0.5) + 0.15 \cdot (1.2 \cdot \#z_2 + 0.2) + 0.5$
			$\#z_1(2) = 1.3525 \cdot \#z_1 + 0.3525 \cdot \#z_2 + 1.075$

Fig. 5. Extending the example from Fig. 3, we compute the expected benefit for $t \in [0, 1, 2]$. Resulting coefficients are cached.

6.1 Benefit Function

To build a stream summary, we define a benefit function $\mathcal{B}$ that scores an element based on the *present benefit* for current aggregation results, and the *expected benefit* for future results.

Present Benefit. The *present benefit* $\mathcal{B}_{\text{pres}}: \mathcal{S} \rightarrow \mathbb{N}$ (for COUNT) of an element e in $\mathcal{S}$ is defined by the number of *complete matches* involving e , i.e., sum of local state counters for final states F in $\mathcal{N}$:

$$\mathcal{B}_{\text{pres}}(e) = \sum_{z_i \in F} e_{\#z_i} \quad (1)$$

We write $\mathcal{B}_{\mathcal{A}_{\text{pres}}}(e)$ to denote all complete matches of e within $\mathcal{A}$.

Expected Benefit. The *expected benefit* $\mathcal{B}_{\text{exp}}: \mathcal{S} \times \mathbb{T} \rightarrow \mathbb{R}_{\geq 0}$ quantifies the expected count of complete matches involving e using its active local state counter $\mathcal{A}_{e\#Z}$ (i.e., (partial) matches with e that may lead to further matches) and the remaining time span Δ_τ , in which e can participate in matches.

We assume that the probability distribution over Σ , coined *alphabet probability* and denoted as $P_\Sigma: \Sigma \rightarrow [0, 1]$, is known. That is, $P_\Sigma(c)$ represents the probability that character $c \in \Sigma$ is the character of the subsequent stream element e . The time window q_τ determines the maximum time an element can lead to (partial) matches. As such, to compute the expected benefit, we induce, for each time step $t \in [0, q_\tau]$, all possible characters $c \in \Sigma$ and measure their influence on the active state counters, weighted by their occurrence probability $P_\Sigma(c)$.

Based on the active local state counters $\#z_i$ in $\mathcal{A}_{e\#Z}$ (denoted here as $\#z_i$ instead of $\mathcal{A}_{\#z_i}$ to ease the presentation), we compute how each $\#z_i$ is affected per time step. For $\#z_i$, the resulting count $\#z_i(t)$ represents the averaged count over all possible words Σ^t weighted by $P_\Sigma(c)$, i.e., the expected match count after t time steps.

Consider Fig. 5, with the regular expression $q_\gamma = A(B^*C)^*D$, the corresponding COUNT rules, and their influences on state counters. Each rule models how a state counter is *influenced*. Since we create a COUNT rule for each transition in $\mathcal{N}$, the number of influences equals the number of transitions $T = \{(z, x, z') \mid z \in Z, x \in \Sigma, z' \in \delta(z, x)\}$. For each active state counter $\#z_i$ in $\mathcal{A}_{e\#Z}$, we determine the *total influence* on $\#z_i$ for a single time step. That is, we determine all influences on $\#z_i$ that correspond to all transitions leading to z_i , i.e., $T_{z_i} = \{(z, x) \mid z \in Z, x \in \Sigma, z_i \in \delta(z, x)\}$. Each tuple

$(z, x) \in T_{z_i}$ represents an influence on z_i . For the automaton in Fig. 3, $T_{z_1} = \{(z_0, A), (z_1, C), (z_2, C)\}$, yielding the influences 1, $\#z_1$, and $\#z_2$, respectively, on $\#z_1$, which results in the total influence $1 + \#z_1 + \#z_2$. Also, each influence is weighted by probability $P_{\Sigma}(c)$ of character c occurring, e.g., for $\#z_1$, we get $1 \cdot P(A) + (\#z_1 + \#z_2) \cdot P(C)$.

For each transition in T , we define an *influence function*, $I_k: \mathbb{T} \rightarrow \mathbb{R}_{\geq 0}$, with $1 \leq k \leq |T|$, indicating the weighted influence of a tuple $(z, x, z') \in T$ at time t , yielding the set of all influences $\mathcal{I} = \{I_1, \dots, I_{|T|}\}$. For every active state counter $\#z_i$ in $\mathcal{A}_{\#Z}$, we define an *influence set* $S_i \subseteq \{1, \dots, |T|\}$, which contains the indices of influence functions that affect $\#z_i$. In Fig. 5, $\#z_1$'s influence set is determined by $S_1 = \{1, 4, 5\}$, indicating that I_1, I_4 , and I_5 affect $\#z_1$.

The computation at time $t + 1$ relies on the results at time t . As such, we derive a linear recurrence relation that operates on the active local state counters $\#z_i$ in $\mathcal{A}_{e_{\#Z}}$ of an element e in $\mathcal{S}$ to compute e 's expected benefit and, therefore, the overall *benefit*.

Specifically, at time $t = 0$, we adopt $\#z_i(0) = \#z_i$ as the initial relation, signifying $\#z_i$'s current value. Then, the general linear recurrence relation for each $t + 1$ and counter $\#z_i$ is defined as:

$$\#z_i(t + 1) = \#z_i(t) + \sum_{k \in S_i} I_k(t) \quad (2)$$

To update $\#z_i(t + 1)$, we add to its previous value, $\#z_i(t)$, the sum of weighted influences on $\#z_i$ at time t . In Fig. 5, at $t = 2$, $\#z_1(2)$ is updated to $1.3525 \cdot \#z_1 + 0.3525 \cdot \#z_2 + 1.075$, that is, $\#z_1(2) = \#z_1(1) + I_1(1) + I_4(1) + I_5(1) = \#z_1(1) + 1 \cdot P(A) + (\#z_1(1) + \#z_2(1)) \cdot P(C)$.

The output of Eq. 2 relies on the initial active local state counters values $\#z_i$ in $\mathcal{A}_{e_{\#Z}}$ of e , which only increase. Hence, the result also captures e 's active present benefit. Thus, to define the *expected benefit function* $\mathcal{B}_{\text{exp}}: \mathcal{S} \times \mathbb{T} \rightarrow \mathbb{R}$, we subtract $\mathcal{B}_{\mathcal{A}_{\text{pres}}}(e)$, estimating the expected count of complete matches e leads to:

$$\mathcal{B}_{\text{exp}}(e, t) = \sum_{z_i \in F} \mathcal{A}_{e_{\#z_i}}(t) - \mathcal{B}_{\mathcal{A}_{\text{pres}}}(e) \quad (3)$$

For instance, let $\mathcal{A}_{e_{\#Z}} = (0, 5, 10, 15)$ in Fig. 5; at $t = 0$, $\mathcal{A}_{e_{\#Z}}$ remains unchanged, while at $t = 1$, $\mathcal{A}_{e_{\#Z}}$ updates to $(0, 7.75, 13, 15.75)$.

Caching Coefficients. The computational cost for $\mathcal{B}_{\text{exp}}(e, t)$ grows as the time window q_{τ} increases. We, therefore, employ a cache that is built during pre-processing using the inputs P_{Σ} and $\mathcal{S}$. Using Fig. 5 for illustration, we note that for each time step t and each active state counter $\#z_i$, a linear combination with constant coefficients is derived. A growing time window q_{τ} modifies only these coefficients, leaving the active state counter $\#z_i$ value anchored to its initial. Hence, for every $t \in [0, q_{\tau}]$, we store the coefficients corresponding to each state counter $\#z_i$. In Fig. 5, for instance, for $\#z_1(2)$, the cached coefficients are 1.3525, 0.3525, and 1.075.

In total, for every $t \in [0, q_{\tau}]$, we keep track of all $O(|Z|)$ state counters by caching at most $O(|Z|)$ coefficients, leading to the memory requirement $O(q_{\tau} \cdot |Z| \cdot |Z|)$, simplifying to $O(q_{\tau} \cdot |Z|^2)$. After the cache is computed, for a given summary element e and time instance t , we can compute $\mathcal{B}_{\text{exp}}(e, t)$ in time $O(|Z|)$ since we have to factor in all cached coefficients for a given input tuple.

6.2 Selection Strategy

The *selection strategy*, denoted as $\psi(\mathcal{S}, e)$, derives a new summary from the current one $\mathcal{S}$ upon the arrival of a stream element e . It may insert e , replace e_{old} in $\mathcal{S}$ by e , or discard e (here, $\mathcal{S}|e_{\text{old}}$ denotes the removal of element e_{old} from $\mathcal{S}$)

$$\psi(\mathcal{S}, e) = \begin{cases} \mathcal{S} \oplus e & \text{if } e \text{ is inserted into } \mathcal{S} \\ (\mathcal{S}|e_{\text{old}}) \oplus e & \text{if } e_{\text{old}} \text{ is replaced by } e \\ \mathcal{S} & \text{if } e \text{ is discarded} \end{cases} \quad (4)$$

The decision of ψ relies on the benefit function $\mathcal{B}$. If the summary capacity has not been reached, a new element e is inserted. If the summary is filled already, the choice between replacement and discarding is done by iterating through the summary from the oldest to the newest element, identifying the element e' with the lowest benefit. For elements outside the active time window $\mathcal{A}$, it holds that $\mathcal{B}(e', 0) = \mathcal{B}_{\text{pres}}(e')$ is their benefit. Upon reaching $\mathcal{A}_{\text{init}}$, we set $e_{\text{old}} = e_{\text{new}} = \mathcal{A}_{\text{init}}$, marking both the oldest and initial newest initiators of $\mathcal{A}$. Note that e_{old} remains constant. As we proceed, encountering a new initiator updates e_{new} . We then determine values $\Delta_{\tau_{\text{old}}} = q_{\tau} - (t - \lambda' [i_{\text{old}}])$ and $\Delta_{\tau_{\text{new}}} = q_{\tau} - (t - \lambda' [i_{\text{new}}])$ with t as the timestamp of e , and $i_{\text{old}}, i_{\text{new}}$ as the indices of $e_{\text{old}}, e_{\text{new}}$. The value $\Delta_{\tau_{\text{old}}}$ specifies the duration during which both e and i_{old} will be involved in matches. Conversely, $\Delta_{\tau_{\text{new}}}$ designates the overall period in which e will participate in matches at most. Based thereon, we define the *benefit function* $\mathcal{B}: \mathcal{S} \times \mathbb{T} \times \mathbb{T} \rightarrow \mathbb{R}_{\geq 0}$:

$$\mathcal{B}(e, \Delta_{\tau_{\text{old}}}, \Delta_{\tau_{\text{new}}}) = \mathcal{B}_{\text{pres}}(e) + \frac{\mathcal{B}_{\text{exp}}(e, \Delta_{\tau_{\text{old}}}) + \mathcal{B}_{\text{exp}}(e, \Delta_{\tau_{\text{new}}})}{2} \quad (5)$$

The function offers a balanced measure of the benefit of e , averaging its potential benefits over the periods $\Delta_{\tau_{\text{old}}}$ and $\Delta_{\tau_{\text{new}}}$.

The idea is based on the following observation: The value of $\Delta_{\tau_{\text{old}}}$ can be small, indicating that e might not contribute to matches for long; however, $\Delta_{\tau_{\text{new}}}$ dictates how long e can lead to matches. On the other hand, relying solely on $\Delta_{\tau_{\text{new}}}$ might not capture the expected benefit accurately, given that e aligns with i_{old} in many (partial) matches, but not for the span of $\Delta_{\tau_{\text{new}}}$.

Finally, we compute $e_{\#Z}$ and $\mathcal{A}_{e_{\#Z}}$ for e , to estimate its benefit $\mathcal{B}(e, \Delta_{\tau_{\text{old}}}, \Delta_{\tau_{\text{new}}})$. If the benefit of e' is smaller, we replace e' with e ; otherwise, e is discarded. Since we compute $\mathcal{B}_{\text{exp}}$ in time $O(|Z|)$ for each element in $\mathcal{S}$, the run time complexity of the selection strategy ψ is given by $O(|\mathcal{S}| \cdot |Z|)$ using $O(q_{\tau} \cdot |Z|^2)$ space.

6.3 Modifications for Other Aggregate Functions

When adopting the SUM aggregate function, the benefit function needs to be adapted. The function for the present benefit is defined analogously, i.e., summing up all values of complete matches. Turning to the expected benefit function, we note that for COUNT, we estimate the count of matches that element e may generate. For SUM, each of these future matches carries a value v . Therefore, the expected benefit for e is augmented by its current present benefit (for SUM), combined with the expected match count of e weighted by v . While this function underestimates the expected SUM value for e , the underestimation is consistent across all summary elements.

7 Experimental Evaluation

To evaluate SuSE, we conducted experiments using the setup described in §7.1. Specifically, we present results on the overall effectiveness, including an ablation and recall study (§7.2); on the sensitivity of the approach (§7.3); on a comparison against state-of-the-art engines for regular expressions and complex event processing (§7.4); and on two real-world datasets (§7.5 and §7.6).

7.1 Experimental Setup

Our implementation and experimental setup is publicly available ¹.

Datasets and Parameters. We used two real-world datasets and synthetic data to assess SuSE regarding the COUNT aggregate. First, we used one month of the Citi Bike dataset [13] that captures information about bike rentals. It contains ≈ 3.8 million events, each of them describing the duration of a trip, the start and end station, a bike ID, and information about the driver.

¹<https://github.com/spurtzel/SuSe>

Our second real-world dataset, NASDAQ [40], contains stock trading data for 462028 events on a minute-by-minute basis. Each entry lists the stock symbol, opening, closing, highest and lowest prices, and the trading volume for that minute.

In synthetic data experiments, we assessed the influence of different parameters on SuSE, i.e., the summary size $|S|$, stream size, number of evaluation timestamps $|\mathcal{E}|$, and the time window q_τ and length of q_γ of a query. To generate the data streams, characters were drawn from alphabets based on different distributions: Zipfian, normal, and uniform. The evaluation timestamps, see Problem 1, were sourced from a Poisson process and a uniform distribution. Yet, the differences between both models have been negligible.

Baselines. We evaluated the performance of our approach against two baselines: A random baseline replaces a randomly chosen element in the summary once its capacity is met. The FIFO baseline consistently replaces the oldest element once its capacity is met.

Our ablation study used a restricted SuSE model, where the benefit function $\mathcal{B}$ is solely based on the present benefit $\mathcal{B}_{\text{pres}}$.

We examined STATESUMMARY's efficiency by comparing it to the state-of-the-art RegEx engine REmatch [51] and CEP engines FlinkCEP [20] and CORE [10]. For REmatch and CORE, we conducted experiments without outputting matches to minimize overhead and ensure fairness. In our Flink setup, while match printing was required to measure processing time for implementation reasons, subsequent evaluations confirmed a negligible impact on the results.

Metrics. We assess the evaluation quality using the *relative recall improvement*. At each evaluation time point $t \in \mathcal{E}$, we divide the current number of complete matches from SuSE by the current number of complete matches from a corresponding baseline. Averaging the sum of these ratios denotes the relative recall improvement.

We also evaluate the *absolute recall*. At each evaluation timestamp, it is calculated by dividing the match count from SuSE by the ground truth, obtained by setting the summary size $|S|$ to the stream size (limited to 2000 in these experiments). Hence, the *absolute recall* coincides with the loss function l_{COUNT} in Problem 1, as both measure how close x (SuSE) is to x_{opt} (the ground truth) and thus yield the same optimum solution, $x = x_{\text{opt}}$. We also examine the *detected matches recall*, i.e., the ratio of complete matches that were present in SuSE to *all* potential matches over time.

We further report performance metrics, such as *throughput* (avg number of elements processed per second), *latency* (avg processing duration per element), and *memory usage* (max resident set size).

Implementation. Our experiments were conducted using a C++ implementation, with 128-bit counters. They ran on a server with 4 Intel Xeon E7-4880 (60 cores, 1TB RAM), except for the experiments shown in Fig. 9, which ran on a system with an Intel i7-1265U processor (10 cores; @4.80 GHz) and 32GB RAM.

7.2 Overall Effectiveness

We first compare against a random baseline (SuSE/Random - S/R) and a FIFO baseline (SuSE/FIFO - S/F). In Fig. 6a - 6b, we assessed how variations in the summary size $|S|$ (100 to 5000), time window size q_τ (10 to 250), and alphabet probability distribution P_Σ (Zipfian, uniform (u), normal (n)) affect the relative recall improvement, using a query with $q_\gamma = a(b^*c)^*d(e|f)g^*$.

SuSE outperforms its baselines by choosing summaries that generate up to 10^{20} more matches. There is no parameter combination, in which SuSE performs worse than a baseline (black dashed line). A clear trend emerges for both plots: the larger q_τ and the smaller $|S|$, the more significant the advantage. As the summary size increases, the baselines find more matches by chance; nevertheless,

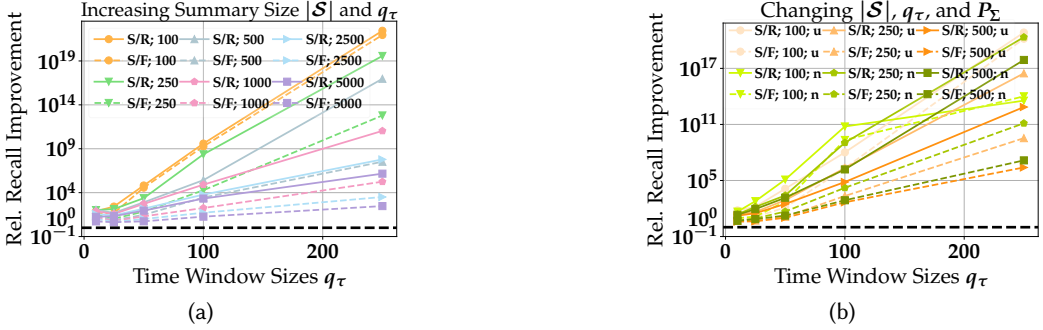


Fig. 6. Examining the impact on relative recall improvement (higher is better) by varying (a) summary and time window sizes and (b) the alphabet probability distribution P_Σ .

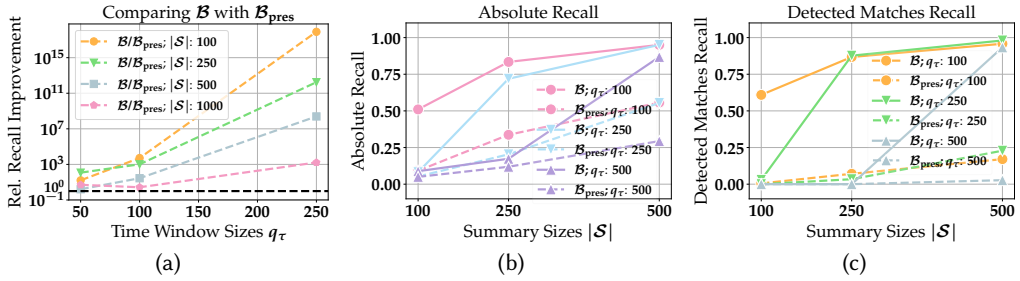


Fig. 7. Examining the impact of the benefit function $\mathcal{B}$ compared to $\mathcal{B}_{\text{pres}}$ (a) on the relative recall improvement, (b) proximity to lower bound in absolute recall, and (c) proximity to lower bound in detected matches recall (higher is better in all plots).

for these parameters, SuSE chooses stream subsequences that generate at least three orders of magnitude more matches on average.

Ablation Study. We compared SuSE's benefit function $\mathcal{B}$ against a setup using solely the present benefit $\mathcal{B}_{\text{pres}}$, i.e., denoted by $\mathcal{B}/\mathcal{B}_{\text{pres}}$, in terms of the relative recall improvement, the related absolute recall, and the detected matches recall. In Fig. 7a, a trend similar to that in Fig. 6a is observed: smaller $|S|$ and larger q_τ lead to better subsequence selections. This is attributed to $\mathcal{B}$ providing a more accurate assessment of an element's future potential. For $q_\tau = 250$, SuSE with $\mathcal{B}$ obtains at least three orders of magnitude more matches on average over solely using $\mathcal{B}_{\text{pres}}$ and up to 10^{18} more for smaller values of $|S|$, showing the effectiveness of $\mathcal{B}$.

Absolute Recall. In Fig. 7b, we evaluated the loss function l_{COUNT} from Problem 1, showing how closely SuSE approaches the optimal recall value.

Here, increasing $|S|$ positively affects recall. This is attributed to the reduced difference between the ground truth's summary size (equal to the stream size) and SuSE's summary size. SuSE clearly attains higher recall values when employing expected benefits $\mathcal{B}$, compared to using solely $\mathcal{B}_{\text{pres}}$. Also, when $|S| \geq q_\tau$, with $\mathcal{B}$, SuSE selects substantially better subsequences, increasing recall by up to 80% – 95%. This trend can be attributed to our assumption in §6.2 regarding the expected benefit $\mathcal{B}$, where we implicitly assumed that at least one time window size of elements fits into the summary. For values $|S| \geq q_\tau$, this assumption is valid, leading to more accurate estimates and an enhanced recall.

Detected Matches Recall. For Fig. 7c, we examined how many of *all* possible matches were present in SuSE, again comparing $\mathcal{B}$ and $\mathcal{B}_{\text{pres}}$. The trends align closely with those observed for the absolute recall. SuSE with $\mathcal{B}$ yields a recall of 90% – 98% for $|S| = 500$, indicating that SuSE chooses rich stream subsequences.

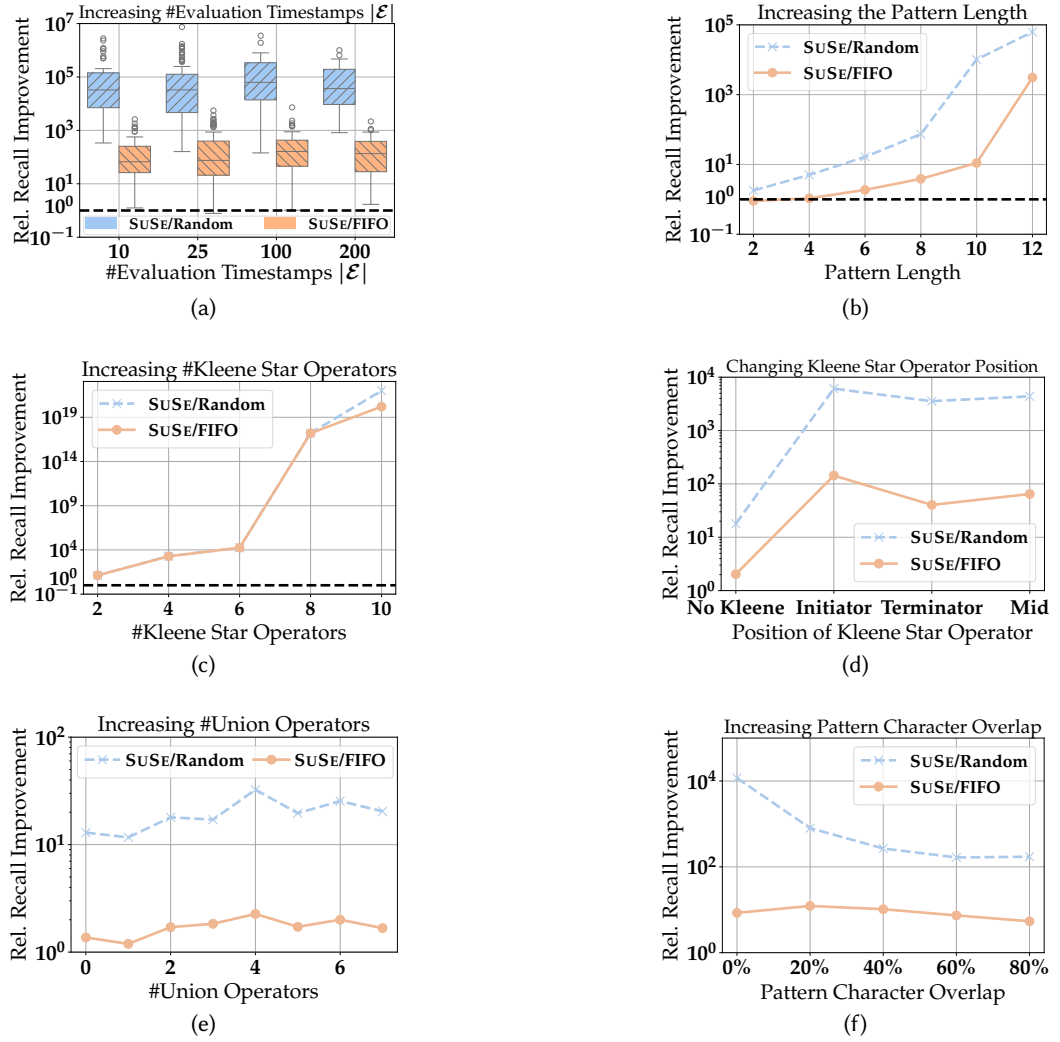


Fig. 8. Examining the impact on relative recall improvement (higher is better) by varying (a) the number of evaluation timestamps $|\mathcal{E}|$, (b) the pattern length, (c) the count of Kleene stars, (d) Kleene star position, (e) union operator count, and (f) query character overlap.

7.3 Sensitivity Analysis

Next, we assess the effect of various parameters on the relative recall improvement. We fix $q_r = 100$ and $|\mathcal{S}| = 500$, and vary stream sizes (10000 - 25000), the numbers of evaluations (25 - 50), and P_{Σ} (Zipfian or uniform). We report averages of over 50 runs.

Number of Evaluation Timestamps. Fig. 8a examines the influence of increasing the number of evaluation timestamps $|\mathcal{E}|$ on the relative recall improvement. As the boxplots indicate, increasing the number of evaluations does not impact the results.

Pattern Length. Fig. 8b shows that as the pattern length grows, selecting subsequences that yield matches becomes increasingly challenging. SuSe addresses this by quantifying the importance of individual elements, prioritizing, for instance, rare elements that are essential to keep to obtain

matches. While for pattern length two, FIFO performs slightly better, for pattern lengths above ten, SuSE is superior by up to almost five orders of magnitude.

Kleene Operator. In Fig. 8c, we employed a regular expression of length twelve and incrementally increased the number of Kleene operators, distributing them randomly. As SuSE is aware of characters with Kleene operators (by COUNT/SUM rules), it selects better summaries.

In Fig. 8d, we examined four regular expressions: one without a Kleene operator and three with the Kleene's position varied. Without Kleene, SuSE's advantages are minimal, while there is a large difference for the other cases, independent of the Kleene position.

Union Operator. We also explored the effect of the union operator in Fig. 8e. We varied the number of unions from zero to seven and generated and combined random sequences for q_Y 's with union operators. Yet, the impact is negligible.

In Fig. 8f, we fixed a query of length ten and progressively increased the character overlap. We randomly selected two characters of the given query for each step, linking them to the query and a union. While the baselines show different trends, in all configurations, SuSE yields significant improvements.

7.4 Efficiency

STATESUMMARY State-of-the-art Comparison. We compared STATESUMMARY's efficiency against the state-of-the-art RegEx engine REmatch and two CEP engines, FlinkCEP and CORE, in Fig. 9. We summarize the idea of the experiment as follows: To assess the current state, i.e., the current number of matches and involvements of individual events, we would need to use a RegEx/CEP engine. However, such engines implement a traditional two-step approach: first building the state, then evaluating aggregate functions over it. Hence, these methods deliver matches and exact results. In contrast, our approach avoids the direct maintenance of (partial) matches and evaluates aggregate functions without match materialization. When the summary is smaller than the word to process, the resulting aggregate value is approximate; however, if the summary is greater than or equal to the word size, the results remain exact.

Given the inferior performance of other RegEx engines (see [51]) and their limitation in computing *all* subsequence matches, we focus on REmatch as a representative RegEx engine. FlinkCEP is an industry-strength tool offering a robust framework for fault-tolerant and distributed stream processing. As such, it offers features that are not implemented in the other engines, which induces a non-negligible overhead during query evaluation that needs to be considered in any comparative analysis. We note for FlinkCEP, though, that parallelizing the workload is not straightforward, since the input stream must be keyed and partitioned while ensuring that *all* matches are found, which is further complicated by operators like Kleene star. While FlinkCEP employs a traditional automata-based evaluation, CORE employs compact match representations, avoiding exponential state growth, and provides match enumeration with output-linear delay.

The evaluated procedures utilized one to two cores on average (i.e., FlinkCEP did not parallelize processing as the input stream is not keyed). The maximum resident set size was ≈ 500 MB for FlinkCEP (increasing to over 1 GB for $x = 256$, and not terminating after several hours), ≈ 409 MB for CORE, ≈ 3.3 MB for REmatch, and ≈ 81 MB for SuSE.

We set the summary size to the corresponding x -value, ensuring that the processed word fits entirely within the summary. Thus, we focused exclusively on STATESUMMARY operations without employing summary selection, computing exact values instead of approximations. We set the time window size such that all elements fit into a single window, maximizing the number of matches to be detected. Also, we ensured that all methods produced the same result when evaluating the aggregate function. We tested with $q_Y = ABCD$ (the REmatch syntax being $!x\{A\}.*!y\{B\}.*!z\{C\}.*!w\{D\}$) for

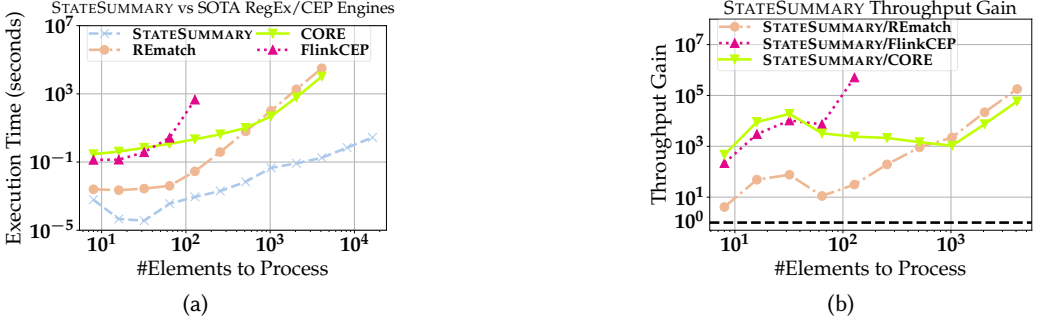


Fig. 9. Comparison against the state of the art, focusing on (a) execution time (lower is better) and (b) throughput gain (higher is better), varying the number of processed elements.

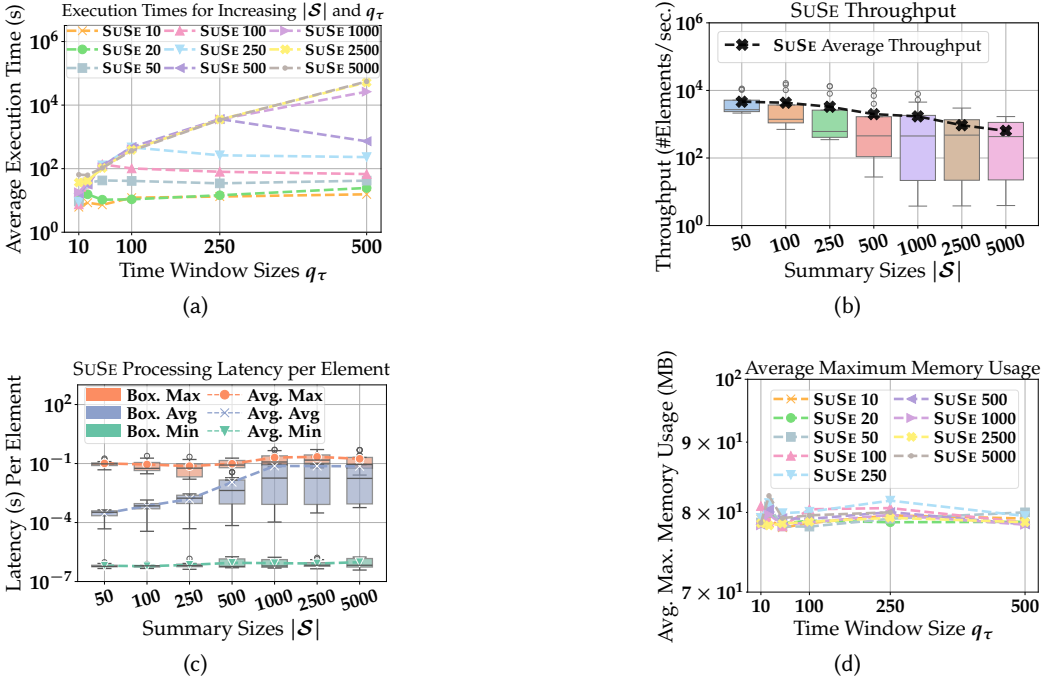


Fig. 10. (a) Execution time (lower is better) for processing 10^5 elements across different summary and time window sizes, (b) throughput (higher is better), (c) average minimum, average maximum, and average processing latency per element (lower is better), and (d) average maximum memory usage (lower is better).

the required processing time of input words of varying sizes. As input, we choose words of the form $A^i B^j C^k D^l$ with $i \in \mathbb{N}^+$, e.g., $A^8 B^8 C^8 D^8$, ensuring an increase in matches for rising x -values.

STATESUMMARY consistently processes words at least an order of magnitude faster than REmatch, and at least three orders of magnitude faster than CORE and Flink. Particularly for words longer than 512 characters, the performance gap between STATESUMMARY and both REmatch and CORE significantly widens, while Flink did not terminate for these values due to overload.

The performance differences induce throughput gains, as shown in Fig. 9b. The ratio indicates how much more quickly STATESUMMARY processed a word than REmatch, CORE, or Flink. The throughput gain always exceeds one, indicating STATESUMMARY's superiority; for longer words, this gain increases significantly.

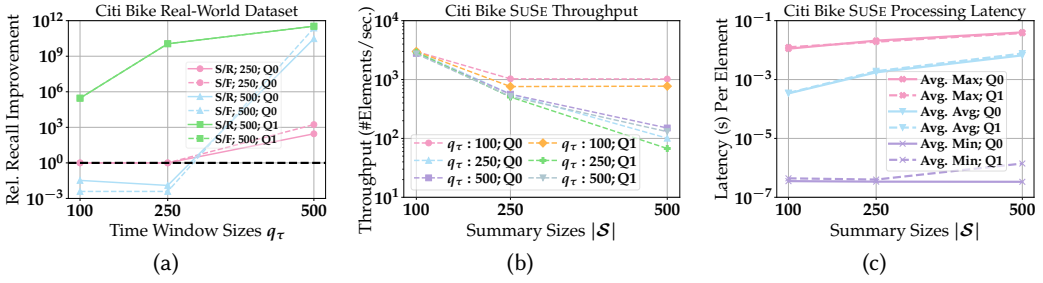


Fig. 11. Citi Bike: (a) relative recall improvement (higher is better), (b) throughput (higher is better), and (c) latency (lower is better).

SuSE's Overall Efficiency. In Fig. 10, we examined how varying $|S|$ (see legend) and q_τ affects the execution time, throughput, processing latency per element, and memory usage of SuSE (employing summary selection) for processing a stream of 10^5 elements. Fig. 10a shows that the larger $|S|$ and the larger q_τ , the longer the execution time. Also, a time window size exists for all procedures up to a summary size of 500, where execution time stabilizes or slightly decreases. The reasoning being that the REMOVE operation is the most expensive one (see §5.2), with cost $\min(q_\tau, |S|)$. Once either $|S| > q_\tau$ or $q_\tau > |S|$, the smaller value of the two parameters becomes the limiting factor for runtime, rendering the influence of the other parameters marginal. With larger q_τ , SuSE chooses better subsequences, which reduces the number of REMOVE operations and explains the drop in runtime for $q_\tau \geq |S|$.

In Fig. 10b, we examined the resulting throughput values for a fixed $q_\tau = 500$. For the same reasons as discussed above, there is a decline in throughput as the summary size expands.

In Fig. 10c, we investigated the resulting average min, max, and average processing latency per element per second. A boxplot shows the latencies for increasing q_τ , while the corresponding line plots denote the averages. A larger $|S|$ leads to elevated processing latencies. However, the median of the average processing latencies consistently remains below 10^{-1} s, showing streaming feasibility.

Finally, the induced memory footprint turned out to be negligible and constant for various $|S|$ and q_τ , see Fig. 10d.

7.5 Case Study: Citi Bike

Character Types and Queries. Fig. 11 denotes our results for the Citi Bike real-world dataset. The character types encompass rides on frequented routes, pinpoint brief rides at busy stations, rides at central stations, and member rides at quieter stations, giving a thorough understanding of the Citi Bike dynamics. We defined a query with $Q0_\gamma = E(H|K)^*E$, recognizing patterns where a ride starts at a busy station, followed by a combination of rides at central and quieter stations, ending with a ride at a busy station, suggesting probable areas for station enhancements or upkeep. The second query, $Q1_\gamma = (E|B|F)^+CEHF$, captures sequences of rides that start at busy stations, popular routes, or extended rides at central spots, followed by short and consecutive rides at busy stations, a trip at a central location, and concluding with a prolonged ride there.

Effectiveness. To examine the relative recall improvement, we varied $|S|$ and q_τ in Fig. 11a. For both queries, not all summary sizes yielded matches. Generally, there is an upward trend in the results with a larger q_τ ; SuSE identifies subsequences that are between three and twelve orders of magnitude superior for $q_\tau = 500$.

Throughput and Latency. Throughput and latency of SuSE are shown in Fig. 11b and Fig. 11c. For smaller values of $|S|$ and q_τ (refer to legend), SuSE's throughput increases. However, with larger values for $|S|$ and q_τ , the lower throughput is attributed to the REMOVE operation. Fig. 11c shows

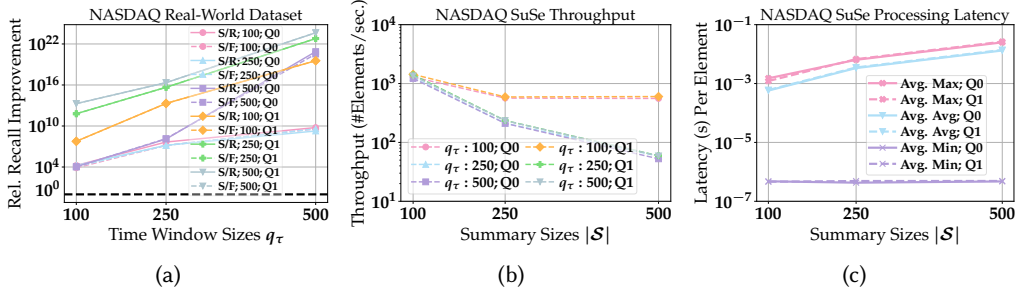


Fig. 12. NASDAQ: (a) relative recall improvement (higher is better), (b) throughput (higher is better), and (c) latency (lower is better).

the avg min, avg max, and avg latency of SuSe, with $q_\tau = 500$. An increase in summary size leads to higher latencies. Yet, the average max latency hovers around 10^{-2} seconds, while the average processing latency remains under 10^{-2} seconds for $|S| = 500$, which meets real-time demands.

7.6 Case Study: NASDAQ

Character Types and Queries. Fig. 12 represents our results for the NASDAQ real-world dataset. We derived character types based on stock market activities, e.g., significant price changes, trade volumes, daily peak or lowest prices, periods of consistent behaviour or fluctuations, and market uncertainty. We formulated a query with $Q0_y = A(B|G)^*A$, detecting an initial price increase, succeeded by any number of price reductions or uncertain market periods, and ending with a subsequent price surge, potentially signalling a fluctuating stock price ascent. The second query, $Q1_y = A^*G(A|B)^*G^*A$, captures zero or more price rises, followed by market uncertainty, zero or more occurrences of either significant price rises or drops, a potential period of market uncertainty, and ends with a price rise.

Effectiveness. In Fig. 12a, we varied $|S|$ and q_τ between 100 – 500 using the $Q0$ and $Q1$ above. Here, for all $|S|$, matches were obtained. Generally, larger $|S|$ and q_τ led SuSe to derive superior stream subsequences, resulting in up to 10^{23} additional matches. However, even at smaller parameter values, SuSe’s subsequences surpassed both baselines by magnitudes ranging from 10^4 to 10^{13} . Notably, for $Q1$, SuSe identified better subsequences due to the increased count of Kleene operators.

Throughput and Latency. In our examination of throughput and processing latency of SuSe, as shown in Fig. 12b and Fig. 12c, a similar trend as for Citi Bike emerges: when decreasing the summary and time window size (see legend), SuSe’s throughput enhances, reaching peaks of up to 1500 elements per second. However, when the $|S|$ and q_τ increase, the throughput decreases to around 70 – 80 elements.

Turning to Fig. 12c, we fixed $q_\tau = 500$ and increased the summary size, resulting in an upward trend of the latencies. Nevertheless, while the average maximum latency approaches 10^{-2} seconds, the average processing latency for $S = 500$ stays just below 10^{-2} seconds, again denoting latencies suitable for real-time computing.

8 Related Work

RegEx Extensions. Practical regular expression matching engines often deviate from the theoretical concept of regular expressions, both in terms of syntax and semantics. Typical syntactical sugar like lookaheads [22], intersection or complement operators [24], or counters [23] allow exponentially more succinct regular expressions. Yet, as a conversion to finite automata is possible, our approach remains applicable. In contrast, the backreference operator [21, 22] to describe unbounded repetitions of substrings is non-regular and leads to an NP-complete matching test [4]. Hence, these

languages are not representable by NFAs, and lifting our approach to these models remains future work.

Regular expressions have been restricted to achieve better (i.e., subquadratic) matching performance, see [8]. Also, so-called deterministic regular expressions have provably better worst-case matching bounds [29]. Since these restrictions represent regular languages, our method is still applicable.

Stream Summaries. Summaries are small data structures that can be updated flexibly and answer certain queries on the input accurately [14]. For instance, a sample of some data may provide an approximate query answer. Sketches are used for the approximate representation of a vector, e.g., in approximate counting or quantile estimation, with Count-Min [15], HyperLogLog [19], GK [27], and KLL [32] being prominent structures. There is also work on summaries for approximate pattern matching under string edit distance [36], where substrings within a streamed text must be found that are within edit distance k of a given pattern.

Our SuSE data structure combines sketching with weighted sampling. The STATESUMMARY provides efficient operations to represent global (and local) state counters, i.e., a vector capturing the current number of (partial) matches per state. Our summary selection algorithm assigns each element a benefit, i.e., a weight, enabling the retention of a bounded sample of stream elements that maximizes the number of matches for a query, and thus approximates the actual number of (partial) matches in the stream.

Subsequences. Subsequences are crucial tools in fields like biological sequence analysis [49] and event stream processing [7, 25, 58]. They are linked to longest common subsequences [9], shortest common supersequences [44], or subsequence counting [17], and are studied in contexts of gap-size [16] and wildcard constraints [34, 35]. However, obtaining aggregated information about subsequence *mappings* within a text has barely been researched. Solely in [17] and [26], dynamic programming algorithms are presented which for a single word y and a text w count the number of mappings $y \leq_m w$. We complement their work using a more holistic approach: We provide aggregated information for a set of patterns derived from a regular expression γ and simultaneously account for partial mappings from $L_{\text{pre}}(\gamma)$. While prior works assume static text, SuSE processes streams, which require online computation.

CEP Optimizations. Queries in CEP use regular expressions for pattern recognition. Here, optimizations include load shedding [11, 53, 59] that discard events or partial matches unlikely to complete; and filters of streams [5] that retain the events that are most likely to lead to complete matches. SuSE also functions as a filter, with its selected summary substream. It resembles these methods in (1) state-based decisions, (2) reducing system load, and (3) minimizing information loss. Yet, SuSE bases decisions on the STATESUMMARY without materializing the state, evading the exponential complexity of automata-based evaluation, and its decisions, unlike [5, 59], are not reliant on a learned model, enhancing streaming practicality.

9 Conclusions

We proposed SuSE, an architecture for regular expression subsequence summarization over data streams. It incorporates a STATESUMMARY data structure that captures a query-specific stream summary through aggregated subsequence match information. We presented a summary selection algorithm, leveraging STATESUMMARY for choosing stream subsequences, which aim to minimize the loss in the aggregated results over time. SuSE is multiple orders of magnitude faster than leading RegEx and CEP engines, while generating aggregates based on richer subsequences than baseline approaches.

References

- [1] 2021. TRE – a Lightweight, Robust, and Efficient POSIX Compliant Regexp Matching Library. <https://github.com/laurikari/tre> Accessed on 2023-10-15.
- [2] Oniguruma – a modern and flexible regular expressions library. 2022. <https://github.com/kkos/oniguruma> Accessed on 2023-10-15.
- [3] Jagrati Agrawal, Yanlei Diao, Daniel Gyllstrom, and Neil Immerman. 2008. Efficient pattern matching over event streams. In *Proceedings of the ACM SIGMOD International Conference on Management of Data, SIGMOD 2008, Vancouver, BC, Canada, June 10-12, 2008*, Jason Tsong-Li Wang (Ed.). ACM, 147–160. doi:10.1145/1376616.1376634
- [4] Alfred V. Aho. 1990. Algorithms for Finding Patterns in Strings. In *Handbook of Theoretical Computer Science, Volume A: Algorithms and Complexity*. 255–300.
- [5] Adar Amir, Ilya Kolchinsky, and Assaf Schuster. 2022. DLACEP: A Deep-Learning Based Framework for Approximate Complex Event Processing. In *Proceedings of the 2022 International Conference on Management of Data*. ACM. doi:10.1145/3514221.3526136
- [6] Anonymized Author(s). 2024. SuSe: Summary Selection for Regular Expression Subsequence Aggregation over Streams. <https://anonymous.4open.science/r/SuSe-8F12/>.
- [7] Alexander Artikis, Alessandro Margara, Martín Ugarte, Stijn Vansummeren, and Matthias Weidlich. 2017. Complex Event Recognition Languages: Tutorial. In *Proceedings of the 11th ACM International Conference on Distributed and Event-based Systems, DEBS 2017, Barcelona, Spain, June 19-23, 2017*. ACM, 7–10. doi:10.1145/3093742.3095106
- [8] Arturs Backurs and Piotr Indyk. 2016. Which Regular Expression Patterns Are Hard to Match?. In *IEEE 57th Annual Symposium on Foundations of Computer Science, FOCS 2016, 9-11 October 2016, Hyatt Regency, New Brunswick, New Jersey, USA*. 457–466. doi:10.1109/FOCS.2016.56
- [9] Ricardo A. Baeza-Yates. 1991. Searching subsequences. *Theoretical Computer Science* 78, 2 (1991), 363–376. doi:10.1016/0304-3975(91)90358-9
- [10] Marco Bucci, Alejandro Grez, Andrés Quintana, Cristian Riveros, and Stijn Vansummeren. 2022. CORE: a complex event recognition engine. *Proceedings of the VLDB Endowment* 15, 9 (May 2022), 1951–1964. doi:10.14778/3538598.3538615
- [11] Koral Chapnik, Ilya Kolchinsky, and Assaf Schuster. 2021. DARLING: Data-Aware Load Shedding in Complex Event Processing Systems. *Proc. VLDB Endow.* 15, 3 (2021), 541–554. <http://www.vldb.org/pvldb/vol15/p541-chapnik.pdf>
- [12] Hao Chen, Yu Chen, and Douglas H. Summerville. 2011. A Survey on the Application of FPGAs for Network Infrastructure Security. *IEEE Communications Surveys & Tutorials* 13, 4 (2011), 541–561. doi:10.1109/surv.2011.072210.00075
- [13] citi Bike. 2022. <http://www.citibikenyc.com/system-data>.
- [14] Graham Cormode. 2022. Current Trends in Data Summaries. *ACM SIGMOD Record* 50, 4 (Jan. 2022), 6–15. doi:10.1145/3516431.3516433
- [15] Graham Cormode and S. Muthukrishnan. 2005. An improved data stream summary: the count-min sketch and its applications. *Journal of Algorithms* 55, 1 (April 2005), 58–75. doi:10.1016/j.jalgor.2003.12.001
- [16] Joel D. Day, Maria Kosche, Florin Manea, and Markus L. Schmid. 2022. Subsequences with Gap Constraints: Complexity Bounds for Matching and Analysis Problems. In *33rd International Symposium on Algorithms and Computation (ISAAC 2022) (LIPIcs, Vol. 248)*, Sang Won Bae and Heejin Park (Eds.). Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl, Germany, 64:1–64:18. doi:10.4230/LIPIcs.ISAAC.2022.64
- [17] Cees Elzinga, Sven Rahmann, and Hui Wang. 2008. Algorithms for subsequence combinatorics. *Theoretical Computer Science* 409, 3 (Dec. 2008), 394–404. doi:10.1016/j.tcs.2008.08.035
- [18] PCRE2 – Perl-Compatible Regular Expressions. 2022. <https://github.com/PCRE2Project/pcr2> Accessed on 2023-10-15.
- [19] Philippe Flajolet, Éric Fusy, Olivier Gandouet, and Frédéric Meunier. 2007. HyperLogLog: the analysis of a near-optimal cardinality estimation algorithm. *Discrete Mathematics & Theoretical Computer Science DMTCS Proceedings vol. AH...*, Proceedings (Jan. 2007). doi:10.46298/dmtcs.3545
- [20] Apache Software Foundation. 2021. Apache Flink. <https://nightlies.apache.org/flink/flink-docs-release-1.14/> Accessed on 2023-10-15.
- [21] Dominik D. Freydenberger and Markus L. Schmid. 2019. Deterministic regular expressions with back-references. *J. Comput. Syst. Sci.* 105 (2019), 1–39. doi:10.1016/J.JCSS.2019.04.001
- [22] Jeffrey E. F. Friedl. 2006. *Mastering regular expressions - understand your data and be more productive: for Perl, PHP, Java, .NET, Ruby, and more (3. ed.)*. O'Reilly. <http://www.oreilly.de/catalog/regex3/index.html>
- [23] Wouter Gelade, Marc Gyssens, and Wim Martens. 2012. Regular Expressions with Counting: Weak versus Strong Determinism. *SIAM J. Comput.* 41, 1 (2012), 160–190. doi:10.1137/100814196
- [24] Wouter Gelade and Frank Neven. 2012. Succinctness of the Complement and Intersection of Regular Expressions. *ACM Trans. Comput. Log.* 13, 1 (2012), 4:1–4:19. doi:10.1145/2071368.2071372

- [25] Nikos Giatrakos, Elias Alevizos, Alexander Artikis, Antonios Deligiannakis, and Minos N. Garofalakis. 2020. Complex event recognition in the Big Data era: a survey. *VLDB J.* 29, 1 (2020), 313–352. doi:10.1007/s00778-019-00557-w
- [26] Ronald I. Greenberg. 2003. Computing the Number of Longest Common Subsequences. arXiv:arXiv:cs/0301034
- [27] Michael Greenwald and Sanjeev Khanna. 2001. Space-efficient online computation of quantile summaries. *ACM SIGMOD Record* 30, 2 (May 2001), 58–66. doi:10.1145/376284.375670
- [28] PCREgrep – A grep program that uses the PCRE regular expression library. 2014. <https://github.com/vmg/pcrere/blob/master/pcregrep.c/> Accessed on 2023-10-15.
- [29] Benoît Groz and Sebastian Maneth. 2017. Efficient testing and matching of deterministic regular expressions. *J. Comput. Syst. Sci.* 89 (2017), 372–399. doi:10.1016/J.JCSS.2017.05.013
- [30] John E. Hopcroft, Rajeev Motwani, and Jeffrey D. Ullman. 2001. Introduction to automata theory, languages, and computation, 2nd edition. *ACM SIGACT News* 32, 1 (March 2001), 60–65. doi:10.1145/568438.568455
- [31] Theodore Johnson, S. Muthukrishnan, and Irina Rozenbaum. 2007. Monitoring Regular Expressions on Out-of-Order Streams. In *2007 IEEE 23rd International Conference on Data Engineering*. IEEE. doi:10.1109/icde.2007.369001
- [32] Zohar Karnin, Kevin Lang, and Edo Liberty. 2016. Optimal Quantile Approximation in Streams. In *2016 IEEE 57th Annual Symposium on Foundations of Computer Science (FOCS)*. IEEE. doi:10.1109/focs.2016.17
- [33] Stephen C. Kleene. 1951. *Representation of events in nerve nets and finite automata*. Technical Report. RAND Corporation, Santa Monica, CA.
- [34] Sarah Kleest-Meißner, Rebecca Sattler, Markus L. Schmid, Nicole Schweikardt, and Matthias Weidlich. 2022. Discovering Event Queries from Traces: Laying Foundations for Subsequence-Queries with Wildcards and Gap-Size Constraints. In *25th International Conference on Database Theory, ICDT 2022 (LIPIcs, Vol. 220)*, Dan Olteanu and Nils Vortmeier (Eds.). Schloss Dagstuhl, 18:1–18:21. doi:10.4230/LIPIcs.ICDT.2022.18
- [35] Sarah Kleest-Meißner, Rebecca Sattler, Markus L. Schmid, Nicole Schweikardt, and Matthias Weidlich. 2023. Discovering Multi-Dimensional Subsequence Queries from Traces - From Theory to Practice. In *Datenbanksysteme für Business, Technologie und Web (BTW 2023), Proceedings (LNI, Vol. P-331)*, Birgitta König-Ries, Stefanie Scherzinger, Wolfgang Lehner, and Gottfried Vossen (Eds.). GI e.V., 511–533. doi:10.18420/BTW2023-24
- [36] Tomasz Kociumaka, Ely Porat, and Tatiana Starikovskaya. 2022. Small-space and streaming pattern matching with k edits. In *2021 IEEE 62nd Annual Symposium on Foundations of Computer Science (FOCS)*. IEEE, 885–896. doi:10.1109/focs52979.2021.00090
- [37] Sailesh Kumar, Balakrishnan Chandrasekaran, Jonathan Turner, and George Varghese. 2007. Curing regular expressions matching algorithms from insomnia, amnesia, and acalculia. In *Proceedings of the 3rd ACM/IEEE Symposium on Architecture for networking and communications systems*. ACM. doi:10.1145/1323548.1323574
- [38] Sailesh Kumar, Sarang Dharmapurikar, Fang Yu, Patrick Crowley, and Jonathan Turner. 2006. Algorithms to accelerate multiple regular expressions matching for deep packet inspection. In *Proceedings of the 2006 conference on Applications, technologies, architectures, and protocols for computer communications*. ACM. doi:10.1145/1159913.1159952
- [39] Boost Regex Library. 2022. <https://github.com/boostorg/regex> Accessed on 2023-10-15.
- [40] NASDAQ Data Link. 2023. <https://data.nasdaq.com/> Accessed: 2023-10-15.
- [41] Alex X. Liu and Eric Norige. 2019. A De-Compositional Approach to Regular Expression Matching for Network Security. *IEEE/ACM Transactions on Networking* 27, 6 (Dec. 2019), 2179–2191. doi:10.1109/tnet.2019.2941920
- [42] Tingwen Liu, Yong Sun, Alex X. Liu, Li Guo, and Binxing Fang. 2012. A Prefiltering Approach to Regular Expression Matching for Network Security Systems. In *Applied Cryptography and Network Security*. Springer Berlin Heidelberg, 363–380. doi:10.1007/978-3-642-31284-7_22
- [43] Lei Ma, Chuan Lei, Olga Poppe, and Elke A. Rundensteiner. 2022. Gloria: Graph-based Sharing Optimizer for Event Trend Aggregation. In *Proceedings of the 2022 International Conference on Management of Data*. ACM. doi:10.1145/3514221.3526145
- [44] David Maier. 1978. The Complexity of Some Problems on Subsequences and Supersequences. *Journal of the ACM* 25, 2 (April 1978), 322–336. doi:10.1145/322063.322075
- [45] Christopher Mutschler and Michael Philippsen. 2014. Adaptive Speculative Processing of Out-of-Order Event Streams. *ACM Trans. Internet Techn.* 14, 1 (2014), 4:1–4:24. doi:10.1145/2633686
- [46] Olga Poppe, Chuan Lei, Lei Ma, Allison Rozet, and Elke A. Rundensteiner. 2021. To Share, or not to Share Online Event Trend Aggregation Over Bursty Event Streams. In *SIGMOD '21: International Conference on Management of Data, Virtual Event, China, June 20-25, 2021*, Guoliang Li, Zhanhui Li, Stratos Idreos, and Divesh Srivastava (Eds.). ACM, 1452–1464. doi:10.1145/3448016.3452785
- [47] Olga Poppe, Chuan Lei, Elke A. Rundensteiner, and David Maier. 2017. GRETA: graph-based real-time event trend aggregation. *Proceedings of the VLDB Endowment* 11, 1 (Sept. 2017), 80–92. doi:10.14778/3151113.3151120
- [48] Olga Poppe, Chuan Lei, Elke A. Rundensteiner, and David Maier. 2019. Event Trend Aggregation Under Rich Event Matching Semantics. In *Proceedings of the 2019 International Conference on Management of Data*. ACM. doi:10.1145/3299869.3319862

- [49] Sven Rahmann. 2006. Subsequence Combinatorics and Applications to Microarray Production, DNA Sequencing and Chaining Algorithms. In *Combinatorial Pattern Matching*. Springer Berlin Heidelberg, 153–164. doi:10.1007/11780441_15
- [50] RE2 regular expression library. 2022. <https://github.com/google/re2> Accessed on 2023-10-15.
- [51] Cristian Riveros, Nicolás Van Sint Jan, and Domagoj Vrgoč. 2023. REmatch: A Novel Regex Engine for Finding All Matches. *Proceedings of the VLDB Endowment* 16, 11 (July 2023), 2792–2804. doi:10.14778/3611479.3611488
- [52] Allison Rozet, Olga Poppe, Chuan Lei, and Elke A. Rundensteiner. 2020. Muse: Multi-query Event Trend Aggregation. In *Proceedings of the 29th ACM International Conference on Information & Knowledge Management*. ACM. doi:10.1145/3340531.3412138
- [53] Ahmad Slo, Sukanya Bhowmik, and Kurt Rothermel. 2020. hSPICE: state-aware event shedding in complex event processing. In *Proceedings of the 14th ACM International Conference on Distributed and Event-based Systems*. ACM. doi:10.1145/3401025.3401742
- [54] Peter Snyder and Chris Kanich. 2015. No Please, After You: Detecting Fraud in Affiliate Marketing Networks.. In *WEIS*.
- [55] Ken Thompson. 1968. Programming Techniques: Regular expression search algorithm. *Commun. ACM* 11, 6 (June 1968), 419–422. doi:10.1145/363347.363387
- [56] Wolfgang Weiss, Víctor Juan Expósito Jiménez, and Herwig Zeiner. 2020. Dynamic Buffer Sizing for Out-of-order Event Compensation for Time-sensitive Applications. *ACM Trans. Sens. Networks* 17, 1 (2020), 1:1–1:23. doi:10.1145/3410403
- [57] Mohamed Zaki and Babis Theodoulidis. 2013. Analyzing Financial Fraud Cases Using a Linguistics-Based Text Mining Approach. *SSRN Electronic Journal* (2013). doi:10.2139/ssrn.2353834
- [58] Haopeng Zhang, Yanlei Diao, and Neil Immerman. 2014. On complexity and optimization of expensive queries in complex event processing. In *International Conference on Management of Data, SIGMOD 2014, Snowbird, UT, USA, June 22-27, 2014*, Curtis E. Dyreson, Feifei Li, and M. Tamer Özsu (Eds.). ACM, 217–228. doi:10.1145/2588555.2593671
- [59] Bo Zhao, Nguyen Quoc Viet Hung, and Matthias Weidlich. 2020. Load Shedding for Complex Event Processing: Input-based and State-based Techniques. In *36th IEEE International Conference on Data Engineering, ICDE 2020, Dallas, TX, USA, April 20-24, 2020*. IEEE, 1093–1104. doi:10.1109/ICDE48307.2020.00099

Received October 2024; revised January 2025; accepted February 2025