

Morpheus Consensus: Excelling on trails and autobahns

Andrew Lewis-Pye   

London School of Economics, UK

Ehud Shapiro  

Weizmann Institute of Science, Israel

Abstract

Recent research in consensus has often focussed on protocols for State-Machine-Replication (SMR) that can handle high throughputs. Such state-of-the-art protocols (generally DAG-based) induce undue overhead when the needed throughput is low, or else exhibit unnecessarily-poor latency and communication complexity during periods of low throughput.

Here we present Morpheus Consensus, which naturally morphs from a quiescent low-throughput leaderless blockchain protocol to a high-throughput leader-based DAG protocol and back, excelling in latency and complexity in both settings. During high-throughput, Morpheus pars with state-of-the-art DAG-based protocols, including Autobahn [15]. During low-throughput, Morpheus exhibits competitive complexity and lower latency than standard protocols such as PBFT [10] and Tendermint [8, 9], which in turn do not perform well during high-throughput.

The key idea of Morpheus is that as long as blocks do not conflict (due to Byzantine behaviour, network delays, or high-throughput simultaneous production) it produces a forkless blockchain, promptly finalizing each block upon arrival. It assigns a leader only if one is needed to resolve conflicts, in a manner and with performance not unlike Autobahn.

2012 ACM Subject Classification Computer systems organization → Dependable and fault-tolerant systems and networks

Keywords and phrases Distributed computing, consensus, quiescence

Digital Object Identifier 10.4230/LIPIcs.OPODIS.2025.35

Related Version *Full Version:* <https://arxiv.org/abs/2502.08465>

1 Introduction

Significant investment in blockchain technology has recently led to renewed interest in research on consensus protocols. Much of this research is focussed on developing protocols that operate efficiently ‘at scale’. In concrete terms, this means looking to design protocols that can handle a high throughput (i.e. high rate of incoming transactions) with low latency (i.e. quick transaction finalization), even when the number of processes (validators) carrying out the protocol is large.

Dealing efficiently with low and high throughput. While blockchains may often need to handle high throughputs, it is not the case that *all* blockchains need to deal with high throughput *all* of the time. For example, various ‘subnets’ or ‘subchains’ may only have to deal with high throughputs infrequently, and should ideally be optimised to deal also with periods of low throughput. The motivation for the present paper therefore stems from a real-world need for consensus protocols that deal efficiently with both high and low throughputs. Specifically, we are interested in a setting where:



© Andrew Lewis-Pye and Ehud Shapiro;

licensed under Creative Commons License CC-BY 4.0

29th International Conference on Principles of Distributed Systems (OPODIS 2025).

Editors: Andrei Arusoaie, Emanuel Onica, Michael Spear, and Sara Tucci-Piergiovanni; Article No. 35;

pp. 35:1–35:20



Leibniz International Proceedings in Informatics

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

1. The processes/validators may be few, but could be up to a few hundred in number.
2. The protocol should be able to handle periods of asynchrony, i.e. should operate efficiently in the partially synchronous setting.
3. The protocol is required to have optimal resilience against Byzantine adversaries, i.e., should be live and consistent so long as less than $1/3$ of processes display Byzantine faults, but should be optimised to deal with the ‘normal case’ that processes are not carrying out Byzantine attacks and that faults are benign (crash or omission failures).
4. There are expected to be some periods of high throughput, meaning that the protocol should ideally match the state-of-the-art during such periods.
5. Often, however, throughput will be low. This means the protocol should also be optimised to give the lowest possible latency during periods of low throughput.
6. Ideally, the protocol should be ‘leaderless’ during periods of low throughput: the use of leaders is to be avoided if possible, since, even without malicious action, leaders who are offline/faulty may cause significant increases in latency.
7. Ideally, the protocol should also be ‘quiescent’, i.e., there should be no need for the sending and storing of new messages when new transactions are not being produced.
8. Transactions may come from *clients* (not belonging to the list of processes/validators), but will generally be produced by the processes themselves.

The main contribution of this paper. We introduce and analyse the Morpheus protocol, which is designed for the setting described above. The protocol is quiescent and has the following properties during periods of low throughput:

- It is leaderless, in the sense that transactions are finalized without the requirement for involvement by leaders.
- Transactions are finalized in time 3δ , where δ is the actual (and unknown) bound on message delays after GST.¹ This more than halves the latency of existing DAG-based protocols and variants such as Autobahn [15] for the low throughput case, and even decreases latency by at least δ when compared with protocols such as PBFT (and even if we suppose leaders for those protocols are non-faulty), since the leaderless property of our protocol negates the need to send transactions to a leader before they can be included in a block.²
- A further advantage over protocols such as PBFT and Tendermint is that crash failures by leaders are not able to impact latency during periods of low throughput.

During periods of high throughput, Morpheus is very similar to Autobahn, and so inherits the benefits of that protocol. In particular:

- It has the same capability to deal with high throughput as DAG-based protocols and variants such as Autobahn, and has the same ability to recover quickly from periods of asynchrony (‘seamless recover’ in the language of Autobahn).
- It has the same latency as Autobahn during high throughput, matching the latency of Sailfish [24], which is the most competitive existing DAG-based protocol in terms of latency.

¹ The partially synchronous setting and associated notions such as GST are formally defined in Section 2.

² See the online version of the paper at <https://arxiv.org/abs/2502.08465> for a detailed analysis of latency and complexity considerations.

84 ■ Morpheus has the same advantages as Autobahn in terms of communication complexity
 85 when compared to DAG-based protocols such as Sailfish, DAG-Rider [17], Cordial Miners
 86 [19], Mysticeti [3] or Shoal [26].

87 Of course, much of the complexity in designing a protocol that operates efficiently in both
 88 low and high throughput settings is to ensure a smooth transition and consistency between
 89 the different modes of operation that the two settings necessitate.

90 **Further contributions of the paper.** In Section 3, we also formalise the task of *Extractable*
 91 *SMR*, as an attempt to make explicit certain implicit assumptions that are often made by
 92 papers in the area. While State-Machine-Replication (SMR) requires correct processes to
 93 finalize *logs* (sequences of transactions) in such a way that *consistency* and *liveness* are
 94 satisfied, it is well understood in the community that some papers describing protocols for
 95 SMR specify protocols that *do not* actually aim to explicitly ensure all correct processes
 96 receive all finalized blocks (required for liveness). Roughly, the protocol instructions suffice
 97 instead to ensure *data availability* (that each finalized block is received by at least one correct
 98 process), and then the protocol is required to establish a total ordering on transactions that
 99 can be extracted via further message exchange, given data availability. Liveness is therefore
 100 only achieved after further message exchange (and via some unspecified method), which
 101 (while a trivial addition if one does not consider communication complexity) is not generally
 102 taken into account when calculating message complexity.

103 In Hotstuff [29], for example, one of the principal aims is to ensure linear message
 104 complexity within *views*. Since this precludes all-to-all communication within views, a
 105 Byzantine leader may finalize a block of transactions in a given view without certain correct
 106 processes even receiving the block. Those correct processes must eventually receive the
 107 block for liveness to be satisfied, but the protocol instructions do not explicitly stipulate
 108 the mechanism by which this should be achieved. While ensuring that all correct processes
 109 receive the block is trivial if one is not concerned with communication complexity (e.g.,
 110 just have each correct process broadcast each finalized block they observe, together with a
 111 quorum certificate verifying that the block is finalized), the messages required to do so are
 112 not counted when analyzing message complexity. The obvious methods of ensuring that the
 113 block is propagated to all correct processes will require more than linear communication
 114 complexity, which undermines the very point of the Hotstuff protocol.

115 The question arises, “what precisely is the task being achieved by such protocols if they
 116 do not satisfy liveness without further message exchange (and so actually fail to achieve the
 117 task of SMR with the communication complexity computed)”. We assert that the task of
 118 Extractable SMR is an appropriate formalisation of the task being achieved, and hope that
 119 the introduction of this notion is a contribution of independent interest.

120 **The structure of the paper.** The paper structure is as follows: Section 2 describes the
 121 basic model and definitions; Section 3 formalises the task of Extractable SMR; Section 4 gives
 122 the intuition behind the Morpheus protocol; Section 5 gives the formal specification of the
 123 protocol; Appendix A formally establishes consistency and liveness; Appendix B discusses
 124 related work. See the online version of the paper at <https://arxiv.org/abs/2502.08465>
 125 for a detailed analysis of latency and complexity considerations.

126 2 The setup

127 We consider a set $\Pi = \{p_0, \dots, p_{n-1}\}$ of n processes. Each process p_i is told i as part of its
 128 input. We consider an adaptive adversary, which chooses a set of at most f processes to

corrupt during the execution, where f is the largest integer less than $n/3$. A process that is corrupted by the adversary is referred to as *Byzantine* and may behave arbitrarily, subject to our cryptographic assumptions (stated below). Processes that are not Byzantine are *correct*.

Cryptographic assumptions. Our cryptographic assumptions are standard for papers in distributed computing. Processes communicate by point-to-point authenticated channels. We use a cryptographic signature scheme, a public key infrastructure (PKI) to validate signatures, a threshold signature scheme [6, 23], and a cryptographic hash function H . The threshold signature scheme is used to create a compact signature of m -of- n processes, as in other consensus protocols [30]. In this paper, $m = n - f$ or $m = f + 1$. The size of a threshold signature is $O(\kappa)$, where κ is a security parameter, and does not depend on m or n . We assume a computationally bounded adversary. Following a common standard in distributed computing and for simplicity of presentation (to avoid the analysis of negligible error probabilities), we assume these cryptographic schemes are perfect, i.e., we restrict attention to executions in which the adversary is unable to break these cryptographic schemes. Hash values are thus assumed to be unique.

Message delays. We consider a discrete sequence of timeslots $t \in \mathbb{N}_{\geq 0}$ in the partially synchronous setting: for some known bound Δ and unknown *Global Stabilization Time* (GST), a message sent at time t must arrive by time $\max\{\text{GST}, t\} + \Delta$. The adversary chooses GST and also message delivery times, subject to the constraints already specified. We write δ to denote the actual (unknown) bound on message delays after GST, noting that δ may be significantly less than the known bound Δ .

Clock synchronization. We do not suppose that the clocks of correct processes are synchronized. For the sake of simplicity, however, we do suppose that the clocks of correct processes all proceed in real time, i.e. if $t' > t$ then the local clock of correct p at time t' is $t' - t$ in advance of its value at time t . This assumption is made only for the sake of simplicity, and our arguments are easily adapted to deal with a setting in which there is a known upper bound on the difference between the clock speeds of correct processes after GST. We suppose all correct processes begin the protocol execution before GST. A correct process may begin the protocol execution with its local clock set to any value.

Transactions. Transactions are messages of a distinguished form. For the sake of simplicity, we consider a setup in which each process produces their own transactions, but one could also adapt the presentation to a setup in which transactions are produced by clients who may pass transactions to multiple processes.

3 Extractable SMR

Informal discussion. State-Machine-Replication (SMR) requires correct processes to finalize *logs* (sequences of transactions) in such a way that *consistency* and *liveness* are satisfied. As noted in Section 1, however, for many papers describing protocols for SMR, the explicit instructions of the protocol *do not* actually suffice to ensure liveness without further message exchange, potentially impacting calculations of message complexity and other measures. Roughly, the protocol instructions do not explicitly ensure that all correct processes receive all finalized blocks, but rather ensure *data availability* (that each finalized block is received by at least one correct process), and then the protocol is required to establish a total ordering on transactions that can be extracted via further message exchange, given data availability. Although it is clear that the protocol can be used to solve SMR given some (as yet unspecified) mechanism for message exchange, the protocol itself does not solve SMR.

174 So, what exactly is the task that the protocol solves?

175 **Extractable SMR (formal definition).** If σ and τ are strings, we write $\sigma \subseteq \tau$ to denote
 176 that σ is a prefix of τ . We say σ and τ are *compatible* if $\sigma \subseteq \tau$ or $\tau \subseteq \sigma$. If two strings are
 177 not compatible, they are *incompatible*. If σ is a sequence of transactions, we write $\mathbf{tr} \in \sigma$ to
 178 denote that the transaction $\mathbf{tr}$ belongs to the sequence σ .

179 If $\mathcal{P}$ is a protocol for *extractable SMR*, then it must specify a function $\mathcal{F}$ that maps
 180 any set of messages to a sequence of transactions. Let M^* be the set of all messages that
 181 are received by at least one (potentially Byzantine) process during the execution. For any
 182 timeslot t , let $M(t)$ be the set of all messages that are received by at least one correct process
 183 at a timeslot $\leq t$. We require the following conditions to hold:

184 *Consistency.* For any M_1 and M_2 , if $M_1 \subseteq M_2 \subseteq M^*$, then $\mathcal{F}(M_1) \subseteq \mathcal{F}(M_2)$.

185 *Liveness.* If correct p produces the transaction $\mathbf{tr}$, there must exist t such that $\mathbf{tr} \in \mathcal{F}(M(t))$.

186 Note that consistency suffices to ensure that, for arbitrary $M_1, M_2 \subseteq M^*$, $\mathcal{F}(M_1)$ and
 187 $\mathcal{F}(M_2)$ are compatible. To see this, note that, by consistency, $\mathcal{F}(M_1) \subseteq \mathcal{F}(M_1 \cup M_2)$ and
 188 $\mathcal{F}(M_2) \subseteq \mathcal{F}(M_1 \cup M_2)$.

189 **Converting protocols for Extractable SMR to protocols for SMR.** In this paper, we
 190 focus on the task of Extractable SMR. One way to convert a protocol for Extractable SMR
 191 into a protocol for SMR is to assume the existence of a gossip network, in which each process
 192 has some (appropriately chosen) constant number of neighbors. Using standard results from
 193 graph theory ([5] Chapter 7), one can assume correct processes form a connected component:
 194 this assumption requires classifying some small number of disconnected processes that would
 195 otherwise be correct as Byzantine. If each correct process gossips each ‘relevant’ protocol
 196 message, then all such messages will eventually be received by all correct processes. Overall,
 197 this induces an extra communication cost per message which is only linear in n . Of course,
 198 other approaches are also possible, and in this paper we will remain agnostic as to the precise
 199 process by which SMR is achieved from Extractable SMR.

200 **4 Morpheus: The intuition**

201 In this section, we informally describe the intuition behind the protocol. The protocol may be
 202 described as ‘DAG-based’, in the sense that each block may *point to* more than one previous
 203 block via the use of hash pointers. The blocks *observed* by any block b are b and all those
 204 blocks observed by blocks that b points to. The set of blocks observed by b is denoted by $[b]$.
 205 If neither of b and b' observe each other, then these two blocks are said to *conflict*. Blocks
 206 will be of three kinds: there exists a unique *genesis* block b_g (which observes only itself), and
 207 all other blocks are either *transaction* blocks or *leader* blocks.

208 **The operation during low throughput.** Roughly, by the ‘low throughput mode’, we
 209 mean a setting in which processes produce blocks of transactions infrequently enough that
 210 correct processes agree on the order in which they are received, meaning that transaction
 211 blocks can be finalized individually upon arrival. Our aim is to describe a protocol that
 212 finalizes transaction blocks with low latency in this setting, and without the use of a leader:
 213 the use of leaders is to be avoided if possible, since leaders who are offline/faulty may cause
 214 significant increases in latency. The way in which Morpheus operates in this setting is simple:

- 215 1. Upon having a new transaction block b to issue, a process p_i will send b to all processes.
- 216 2. If they have not seen any blocks conflicting with b , other processes then send a *1-vote* for
 217 b to all processes.

- 218 3. Upon receiving $n - f$ 1-votes for b , and if they still have not seen any block conflicting
 219 with b , each correct process will send a *2-vote* for b to all others.
 220 4. Upon receiving $n - f$ 2-votes for b , a process regards b as finalized.

221 Recall that δ is the actual (unknown) bound on message delays after GST. If the new
 222 transaction block b is created at time $t > \text{GST}$, then the procedure above causes all correct
 223 processes to regard b as finalized by time $t + 3\delta$.

224 Which blocks should a new transaction block b point to? For the sake of concreteness, let
 225 us specify that if there is a *sole tip* amongst the blocks received by p_i , i.e., if there exists a
 226 unique block b' amongst those received by p_i which observes all other blocks received by p_i ,
 227 then p_i should have b point to b' . To integrate with our approach to the ‘high throughput
 228 mode’, we also require that b should point to the last transaction block created by p_i .
 229 Generally, we will only require transaction blocks to point to at most two previous blocks.
 230 This avoids the downside of many DAG-based protocols that all blocks require $O(n)$ pointers
 231 to previous blocks.

232 **Moving to high throughput.** When conflicting transaction blocks are produced, we need
 233 a method for ordering them. The approach we take is to use *leaders*, who produce a second
 234 type of block, called *leader* blocks. These leader blocks are used to specify the required total
 235 ordering.

236 *Views.* In more detail, the instructions for the protocol are divided into *views*, each with a
 237 distinct leader. If a particular view is operating in ‘low throughput’ mode and conflicting
 238 blocks are produced, then some time may pass during which a new transaction block fails to
 239 be finalized. In this case, correct processes will complain, by sending messages indicating
 240 that they wish to move to the next view. Once processes enter the next view, the leader of
 241 that view will then continue to produce leader blocks so long as the protocol remains in high
 242 throughput mode. Each of these leader blocks will point to all *tips* (i.e. all blocks which are
 243 not observed by any others) seen by the leader, and will suffice to specify a total ordering on
 244 the blocks they observe.

245 *The two phases of a view.* Each view is thus of potentially unbounded length and consists
 246 of two phases. During the first phase, the protocol is in high throughput mode, and is
 247 essentially the same as Autobahn.³ Processes produce transaction blocks, each of which
 248 just points to their last produced transaction block. Processes do not send 1 or 2-votes for
 249 transaction blocks during this phase, but rather vote for leader blocks, which, when finalized,
 250 suffice to specify the required total ordering on transactions. Leader blocks are finalized as
 251 in PBFT, after two rounds of voting. If a time is reached after which transaction blocks
 252 arrive infrequently enough that leader blocks are no longer required, then the view enters
 253 a second phase, during which processes vote on transaction blocks and attempt to finalize
 254 them without the use of a leader.

255 **How to produce the total ordering.** For protocols in which each block points to a
 256 single predecessor, the total ordering of transactions specified by a finalized block b is clear:
 257 the ordering on transactions is just that inherited by the sequence of blocks below b and
 258 the transactions they contain. In a context where each block may point to multiple others,
 259 however, we have extra work to do to specify the required total ordering on transactions.

³ We note that Autobahn includes the option of various optimisations (with corresponding tradeoffs) that can be used to further reduce latency in certain ‘good’ scenarios (where all processes act correctly, for example). For the sake of simplicity we do not include these optimisations in our formal description of Morpheus in Section 5, but the online version of this paper discusses those options.

The approach we take is similar to many DAG-based protocols (e.g. [19]). Given any sequence of blocks S , we let $\text{Tr}(S)$ be the corresponding sequence of transactions, i.e. if $b_1, \dots, b_k$ is the subsequence of S consisting of the transaction blocks in S , then $\text{Tr}(S)$ is $b_1.\text{Tr} * b_2.\text{Tr} * \dots * b_k.\text{Tr}$, where $*$ denotes concatenation, and where $b.\text{Tr}$ is the sequence of transactions in b . We suppose given $\tau^\dagger$ such that, for any set of blocks B , $\tau^\dagger(B)$ is a sequence of blocks that contains each block in B precisely once, and which respects the observes relation: if $b, b' \in B$ and b' observes b , then b appears before b' in $\tau^\dagger(B)$. Each transaction/leader block b will contain q which is a 1-Quorum-Certificate (1-QC), i.e., a threshold signature formed from $n - f$ 1-votes, for some previous block: this will be recorded as the value $b.1\text{-QC} = q$, while, if q is a 1-QC for b' , then we set $q.b = b'$. QCs are ordered first by the view of the block to which they correspond, then by the type of the block (leader or transaction, with the latter being greater), and then by the height of the block. We then define $\tau(b)$ by recursion:

- $\tau(b_g) = b_g$.
- If $b \neq b_g$, then let $q = b.1\text{-QC}$ and set $b' = q.b$. Then $\tau(b) = \tau(b') * \tau^\dagger([b] - [b'])$.

Given any set of messages M , let M' be the largest set of blocks in M that is *downward closed*, i.e. such that if $b \in M'$ and b observes b' , then $b' \in M'$. Let q be a maximal 2-QC in M such that $q.b \in M'$, and set $b = q.b$, or if there is no such 2-QC in M , set $b = b_g$. We define $\mathcal{F}(M)$ to be $\text{Tr}(\tau(b))$.

Maintaining consistency. Consistency is formally established in Appendix A, and uses a combination of techniques from PBFT, Tendermint, and previous DAG-based protocols. Roughly, the argument is as follows. When the protocol moves to a new view, consistency will be maintained using the same technique as in PBFT. Upon entering the view, each process sends a ‘new-view’ message to the leader, specifying the greatest 1-QC they have seen. Upon producing a first leader block b for the view, the leader must then justify the choice of $b.1\text{-QC}$ by listing new-view messages signed by $n - f$ distinct processes in Π . The value $b.1\text{-QC}$ must be greater than or equal to all 1-QCs specified in those new-view messages. If any previous block b' has received a 2-QC, then at least $f + 1$ correct processes must have seen a 1-QC for b' , meaning that $b.1\text{-QC}$ must be greater than or equal to that 1-QC. Subsequent leader blocks b'' for the view just set $b''.1\text{-QC}$ to be a 1-QC for the previous leader block.

To maintain consistency between finalized transaction blocks and between leader and transaction blocks within a single view, we also have each transaction block specify q which is 1-QC for some previous block. Correct processes will not vote for the transaction block unless q is greater than or equal to any 1-QC they have previously received.

Overall, the result of these considerations is that, if two blocks b and b' receive 2-QCs q and q' respectively, with q greater than q' , then the iteration specifying $\tau(b)$ (as detailed above) proceeds via b' , so that $\tau(b)$ extends $\tau(b')$.

0-votes. While operating in low throughput, a 1-QC for a block b suffices to ensure both data availability, i.e. that some correct process has received the block, and *non-equivocation*, i.e. two conflicting blocks cannot both receive 1-QCs. When operating in high throughput, however, transaction blocks will not receive 1 or 2-votes. In this context, we still wish to ensure data availability. It is also useful to ensure that each individual process does not produce transaction blocks that conflict with each other, so as to bound the number of tips that may be created. To this end, we make use of *0-votes*, which may be regarded as weaker than standard votes for a block:

1. Upon having a new transaction block b to issue, a process p_i will send b to all processes.

- 306 2. If the block is properly formed, and if other processes have not seen p_i produce any
 307 transaction blocks conflicting with b , then they will send a 0 -vote for b back to p_i . Note
 308 that 0 -votes are sent only to the block creator, rather than to all processes.
- 309 3. Upon receiving $n - f$ 0 -votes for b , p_i will then form a 0 -QC for b and send this to all
 310 processes.

311 When a block b' wishes to point to b , it will include a z -QC for b (for some $z \in \{0, 1, 2\}$). As
 312 a consequence, any process will be able to check that b' is valid/properly formed without
 313 actually receiving the blocks that b' points to: the existence of QCs for those blocks suffices
 314 to ensure that they are properly formed (and that at least one correct process has those
 315 blocks), and other requirements for the validity of b' can be checked by direct inspection. For
 316 this to work, votes (and QCs) must specify certain properties of the block beyond its hash,
 317 such as the height of the block and the block creator. The details are given in Section 5.

318 5 Morpheus: the formal specification

319 The pseudocode uses a number of local variables, functions, objects and procedures, detailed
 320 below. In what follows, we suppose that, when a correct process sends a message to ‘all
 321 processes’, it regards that message as immediately received by itself. All messages are signed
 322 by the sender. For any variable x , we write $x \downarrow$ to denote that x is defined, and $x \uparrow$ to denote
 323 that x is undefined. Table 1 lists all message types.

Message type	Description
Blocks	
Genesis block	Unique block of height 0
Transaction blocks	Contain transactions
Leader blocks	Used to totally order transaction blocks
Votes and QCs	
0 -votes	Guarantee data availability and non-equivocation in high throughput
1 -votes	Sent during 1st round of voting on a block
2 -votes	Sent during 2nd round of voting on a block
z -QC, $z \in \{0, 1, 2\}$	formed from $n - f$ z -votes
View messages	
End-view messages	Indicate wish to enter next view
$(v + 1)$ -certificate	Formed from $f + 1$ end-view v messages
View v message	Sent to the leader at start of view v

■ **Table 1** Message types.

324 **The genesis block.** There exists a unique *genesis block*, denoted b_g . For any block b , $b.type$
 325 specifies the type of the block b , $b.view$ is the *view* corresponding to the block, $b.h$ specifies
 326 the *height* of the block, $b.auth$ is the block creator, and $b.slot$ specifies the *slot* corresponding
 327 to the block. For b_g , we set:

328 ■ $b_g.type = \text{gen}$, $b_g.view = -1$, $b_g.h = 0$, $b_g.auth = \perp$, $b_g.slot = 0$.

A comment on the use of slots. Each block will either be the genesis block, a transaction block, or a leader block. If $p_i \in \Pi$ is correct then, for $s \in \mathbb{N}_{\geq 0}$, p_i will produce a single transaction block b with $b.\text{slot} = s$ before producing any transaction block b' with $b'.\text{slot} = s + 1$. Similarly, if $p_i \in \Pi$ is correct then, for $s \in \mathbb{N}_{\geq 0}$, p_i will produce a single leader block b with $b.\text{slot} = s$ before producing any leader block b' with $b'.\text{slot} = s + 1$.

329

330 **z -votes.** For $z \in \{0, 1, 2\}$, a z -vote for the block b is a message of the form $(z, b.\text{type}, b.\text{view},$
 331 $b.\text{h}, b.\text{auth}, b.\text{slot}, H(b))$, signed by some process in Π . The reason votes include more in-
 332 formation than just the hash of the block is explained in Section 4. A z -quorum for b is
 333 a set of $n - f$ z -votes for b , each signed by a different process in Π . A z -QC for b is the
 334 message $m = (z, b.\text{type}, b.\text{view}, b.\text{h}, b.\text{auth}, b.\text{slot}, H(b))$ together with a threshold signature
 335 for m , formed from a z -quorum for b using the threshold signature scheme.

336 **QCs.** By a QC for the block b , we mean a z -QC for b , for some $z \in \{0, 1, 2\}$. If q is a z -QC for
 337 b , then we set $q.b = b$, $q.z = z$, $q.\text{type} = b.\text{type}$, $q.\text{view} = b.\text{view}$, $q.\text{h} = b.\text{h}$, $q.\text{auth} = b.\text{auth}$,
 338 $q.\text{slot} = b.\text{slot}$. We define a preordering $\leq$ on QCs as follows: QCs are preordered first by
 339 view, then by type with lead $< \text{Tr}$, and then by height.⁴

340 **The variable M_i .** Each process p_i maintains a local variable M_i , which is automatically
 341 updated and specifies the set of all received messages. Initially, M_i contains b_g and a 1-QC
 342 for b_g .

343 **Transaction blocks.** Each transaction block b is entirely specified by the following values:

- 344 ■ $b.\text{type} = \text{Tr}$, $b.\text{view} = v \in \mathbb{N}_{\geq 0}$, $b.\text{h} = h \in \mathbb{N}_{> 0}$, $b.\text{slot} = s \in \mathbb{N}_{\geq 0}$.
- 345 ■ $b.\text{auth} \in \Pi$: the block creator.
- 346 ■ $b.\text{Tr}$: a sequence of transactions.
- 347 ■ $b.\text{prev}$: a non-empty set of QCs for blocks of height $< h$.
- 348 ■ $b.1\text{-QC}$: a 1-QC for a block of height $< h$.

349 If $b.\text{prev}$ contains a QC for b' , then we say that b points to b' . For b to be *valid*, we require
 350 that it is of the form above and:

- 351 1. b is signed by $b.\text{auth}$.
- 352 2. If $s > 0$, b points to b' with $b'.\text{type} = \text{Tr}$, $b'.\text{auth} = b.\text{auth}$ and $b'.\text{slot} = s - 1$.
- 353 3. If b points to b' , then $b'.\text{view} \leq b.\text{view}$.
- 354 4. If $h' = \max\{b'.\text{h} : b \text{ points to } b'\}$, then $h = h' + 1$.

355 We suppose correct processes ignore transaction blocks that are not valid. In what follows we
 356 therefore adopt the convention that, by a ‘transaction block’, we mean a ‘valid transaction
 357 block’.

⁴ For the sake of completeness, if $q.\text{view} = q'.\text{view}$, $q.\text{type} = q'.\text{type}$, and $q.\text{h} = q'.\text{h}$, then we set $q \leq q'$ and $q' \leq q$. We will show that, in this case, $q.b = q'.b$.

A comment on transaction blocks. During periods of high throughput, a transaction block produced by p_i for slot s will just point to p_i 's transaction block for slot $s - 1$. During periods of low throughput, if there is a unique block b' received by p_i that does not conflict with any other block received by p_i , any transaction block b produced by p_i will also point to b' (so that b does not conflict with b').

The use of $b.1\text{-QC}$ is as follows: once correct p_i sees a 1-QC q , it will not vote for any transaction block b unless $b.1\text{-QC}$ is greater than or equal to q . Ultimately, this will be used to argue that consistency is satisfied.

358

When blocks observe each other. The genesis block *observes* only itself. Any other block b observes itself and all those blocks observed by blocks that b points to. If two blocks do not observe each other, then they *conflict*. We write $[b]$ to denote the set of all blocks observed by b .

The leader of view v . The leader of view v , denoted $\text{lead}(v)$, is process p_i , where $i = v \bmod n$.

End-view messages. If process p_i sees insufficient progress during view v , it may send an end-view v message of the form (v) , signed by p_i . By a *quorum* of end-view v messages, we mean a set of $f + 1$ end-view v messages, each signed by a different process in Π . If p_i receives a quorum of end-view v messages before entering view $v + 1$, it will combine them (using the threshold signature scheme) to form a $(v + 1)$ -certificate. Upon first seeing a $(v + 1)$ -certificate, p_i will send this certificate to all processes and enter view $v + 1$. This ensures that, if some correct process is the first to enter view $v + 1$ after GST, all correct processes enter that view (or a later view) within time Δ .

View v messages. When p_i enters view v , it will send to $\text{lead}(v)$ a view v message of the form (v, q) , signed by p_i , where q is a maximal amongst 1-QCs seen by p_i . We say that q is the 1-QC *corresponding* to the view v message (v, q) .

A comment on view v messages. The use of view v messages is to carry out view changes in the same manner as PBFT. When producing the first leader block b of the view, the leader must include a set of $n - f$ view v messages, which act as a *justification* for the block proposal: the value $b.1\text{-QC}$ must be greater than or equal all 1-QCs corresponding to those $n - f$ view v messages. For each subsequent leader block b' produced in the view, $b'.1\text{-QC}$ must be a 1-QC for the previous leader block (i.e., that for the previous slot). The argument for consistency will thus employ some of the same methods as are used to argue consistency for PBFT.

376

Leader blocks. Each *leader block* b is entirely specified by the following values:

- $b.\text{type} = \text{lead}$, $b.\text{view} = v \in \mathbb{N}_{\geq 0}$, $b.h = h \in \mathbb{N}_{> 0}$, $b.\text{slot} = s \in \mathbb{N}_{\geq 0}$.
- $b.\text{auth} \in \Pi$: the block creator.
- $b.\text{prev}$: a non-empty set of QCs for blocks of height $< h$.
- $b.1\text{-QC}$: a 1-QC for a block of height $< h$.
- $b.\text{just}$: a (possibly empty) set of view v messages.

As for transaction blocks, if $b.\text{prev}$ contains a QC for b' , then we say that b *points to* b' . For b to be *valid*, we require that it is of the form described above and:

1. b is signed by $b.\text{auth}$ and $b.\text{auth} = \text{lead}(v)$.
2. If b points to b' , then $b'.\text{view} \leq b.\text{view}$.

386

- 387 3. If $h' = \max\{b'.h : b \text{ points to } b'\}$, then $h = h' + 1$.
- 388 4. If $s > 0$, b points to a unique b^* with $b^*.type = \text{lead}$, $b^*.auth = b.auth$ and $b^*.slot = s - 1$.
- 389 5. If $s = 0$ or $b^*.view < v$, then $b.just$ contains $n - f$ view v messages, each signed by a
390 different process in Π . This set of messages is called a *justification* for the block.
- 391 6. If $s = 0$ or $b^*.view < v$, then $b.1\text{-QC}$ is greater than or equal to all 1-QCs corresponding
392 to view v messages in $b.just$.
- 393 7. If $s > 0$ and $b^*.view = v$, then $b.1\text{-QC}$ is a 1-QC for b^* .

394 As with transaction blocks, we suppose correct processes ignore leader blocks that are not
395 valid. In what follows we therefore adopt the convention that, by a ‘leader block’, we mean a
396 ‘valid leader block’.

A comment on leader blocks. The conditions for validity above are just those required to carry out a PBFT-style approach to view changes (as discussed previously). The first leader block of the view must include a justification for the block proposal (to guarantee consistency). Subsequent leader blocks in the view simply include a 1-QC for the previous leader block (i.e., that for the previous slot).

397
398 **The variable Q_i .** Each process p_i maintains a local variable Q_i , which is automatically
399 updated and, for each $z \in \{0, 1, 2\}$, stores at most one z -QC for each block: For $z \in \{0, 1, 2\}$,
400 if p_i receives⁵ a z -quorum or a z -QC for b , and if Q_i does not contain a z -QC for b , then
401 p_i automatically enumerates a z -QC for b into Q_i (either the z -QC received, or one formed
402 from the z -quorum received).

403 We define the ‘observes’ relation $\succeq$ on Q_i to be the minimal preordering satisfying
404 (transitivity and):

- 405 ■ If $q, q' \in Q_i$, $q.type = q'.type$, $q.auth = q'.auth$ and $q.slot > q'.slot$, then $q \succeq q'$.
- 406 ■ If $q, q' \in Q_i$, $q.type = q'.type$, $q.auth = q'.auth$, $q.slot = q'.slot$, and $q.z \geq q'.z$, then
407 $q \succeq q'$.
- 408 ■ If $q, q' \in Q_i$, $q.b = b$, $q'.b = b'$, $b \in M_i$ and b points to b' , then $q \succeq q'$.

409 We note that the observes relation $\succeq$ depends on Q_i and M_i , and is stronger than the
410 preordering $\geq$ we defined on z -QCs previously, in the following sense: if q and q' are z -QCs
411 with $q \succeq q'$, then $q \geq q'$, while the converse may not hold. When we refer to the ‘greatest’ QC
412 in a given set, or a ‘maximal’ QC in a given set, this is with reference to the $\geq$ preordering,
413 unless explicitly stated otherwise. If $q.type = q'.type$, $q.auth = q'.auth$ and $q.slot = q'.slot$,
414 then it will follow that $q.b = q'.b$.

A comment on the observes relation on Q_i . When p_i receives $q, q' \in Q_i$, it may not be immediately apparent whether $q.b$ observes $q'.b$. The observes relation defined on Q_i above is essentially that part of the observes relation on blocks that p_i can testify to, given the messages it has received (while also distinguishing the ‘level’ of the QC).

415
416 **The tips of Q_i .** The *tips* of Q_i are those $q \in Q_i$ such that there does not exist $q' \in Q_i$ with
417 $q' \succ q$ (i.e. $q' \succeq q$ and $q \not\succeq q'$). The protocol ensures that Q_i never contains more than $2n$
418 tips: The factor 2 here comes from the fact that leader blocks produced by correct p_i need
419 not observe all transaction blocks produced by p_i (and vice versa).

⁵ Here, we include the possibility that p_i receives the z -QC inside a message, such as in $b'.prev$ for a received block b'

420 **Single tips.** We say $q \in Q_i$ is a *single tip* of Q_i if $q \succeq q'$ for all $q' \in Q_i$. We say $b \in M_i$ is a
 421 *single tip* of M_i if there exists q which is a single tip of Q_i and b is the unique block in M_i
 422 pointing to $q.b$.

A comment on single tips. When a transaction block is a single tip of M_i , this will enable p_i to send a 1-vote for the block. Leader blocks do not have to be single tips for correct processes to vote for them.

423

424 **The voted function.** For each $i, j, s, z \in \{0, 1, 2\}$ and $x \in \{\text{lead}, \text{Tr}\}$, the value $\text{voted}_i(z, x, s, p_j)$
 425 is initially 0. When p_i sends a z -vote for a block b with $b.\text{type} = x$, $b.\text{auth} = p_j$, and $b.\text{slot} = s$,
 426 it sets $\text{voted}_i(z, x, s, p_j) := 1$. Once this value is set to 1, p_i will not send a z -vote for any
 427 block b' with $b'.\text{type} = x$, $b'.\text{auth} = p_j$, and $b'.\text{slot} = s$.

428 **The phase during the view.** For each i and v , the value $\text{phase}_i(v)$ is initially 0. Once p_i
 429 votes for a transaction block during view v , it will set $\text{phase}_i(v) := 1$, and will then not vote
 430 for leader blocks within view v .

A comment on the phase during a view. As noted previously, each view can be thought of as consisting of two phases. Initially, the leader is responsible for finalizing transactions. If, after some time, the protocol enters a period of low throughput, then the leader will stop producing leader blocks, and transactions blocks can then be finalized directly. Once a process votes for a transaction block, it may be considered as having entered the low throughput phase of the view. The requirement that it should not then vote for subsequent leader blocks in the view is made so as to ensure consistency between finalized leader blocks and transaction blocks within the view.

431

432 **When blocks are final.** Process p_i regards $q \in Q_i$ (and $q.b$) as *final* if there exists $q' \in Q_i$
 433 such that $q' \succeq q$ and q' is a 2-QC (for any block).

434 **The function $\mathcal{F}$.** This is defined exactly as specified in Section 4.

435 **The variables view_i and $\text{slot}_i(x)$ for $x \in \{\text{lead}, \text{Tr}\}$.** These record the present view and
 436 slot numbers for p_i .

437 **The PayloadReady_i function.** We remain agnostic as to how frequently processes should
 438 produce transaction blocks, i.e. as to whether processes should produce transaction blocks
 439 immediately upon having new transactions to process, or wait until they have a set of new
 440 transactions of at least a certain size. We suppose simply that:

- 441 ■ Extraneous to the explicit instructions of the protocol, PayloadReady_i may be set to 1
 442 at some timeslots of the execution.
- 443 ■ If $\text{PayloadReady}_i = 1$ and $\text{slot}_i(\text{Tr}) = s > 0$, then there exists $q \in Q_i$ with $q.\text{auth} = p_i$,
 444 $q.\text{type} = \text{Tr}$ and $q.\text{slot} = s - 1$.

A comment on the PayloadReady_i function. The second requirement above is required so that p_i can ensure that the new transaction block it forms can point to its transaction block for the previous slot.

445

446 **The procedure MakeTrBlock_i .** When p_i wishes to form a new transaction block b , it will
 447 run this procedure, by executing the following instructions:

- 448 1. Set $b.\text{type} := \text{Tr}$, $b.\text{auth} := p_i$, $b.\text{view} := \text{view}_i$, $b.\text{slot} := \text{slot}_i(\text{Tr})$.

2. Let $s := \text{slot}_i(\text{Tr})$. If $s > 0$, then let $q_1 \in Q_i$ be such that $q_1.\text{auth} = p_i$, $q_1.\text{type} = \text{Tr}$ and $q_1.\text{slot} = s - 1$. If $s = 0$, let q_1 be a 1-QC for b_g . Initially, set $b.\text{prev} := \{q_1\}$.
3. If there exists $q_2 \in Q_i$ which is a single tip of Q_i , then enumerate q_2 into $b.\text{prev}$.
4. If $h' = \max\{q.h : q \in b.\text{prev}\}$, then set $b.h := h' + 1$.
5. Let q be the greatest 1-QC in Q_i . Set $b.1\text{-QC} := q$.
6. Sign b with the values specified above, and send this block to all processes.
7. Set $\text{slot}_i(\text{Tr}) := \text{slot}_i(\text{Tr}) + 1$;

The boolean LeaderReady_i . At any time, this boolean is equal to 1 iff either of the following conditions are satisfied, setting $v = \text{view}_i$:

1. Process p_i has not yet produced a block b with $b.\text{view} = v$ and $b.\text{type} = \text{lead}$, and both:
 - a. Process p_i has received view v messages signed by at least $n - f$ processes in Π .
 - b. $\text{slot}_i(\text{lead}) = 0$ or Q_i contains q with $q.\text{auth} = p_i$, $q.\text{type} = \text{lead}$, $q.\text{slot} = \text{slot}_i(\text{lead}) - 1$.
2. Process p_i has previously produced a block b with $b.\text{view} = v$ and $b.\text{type} = \text{lead}$, and Q_i contains a 1-QC for b' with $b'.\text{auth} = p_i$, $b'.\text{type} = \text{lead}$, $b'.\text{slot} = \text{slot}_i(\text{lead}) - 1$.

A comment on the boolean LeaderReady_i . If p_i is the leader for view v , then before producing the first leader block of the view, it must receive view v messages from $n - f$ different processes, and must also receive a QC for the last leader block it produced (if any). Before producing any subsequent leader block in the view, it must receive a 1-QC for the previous leader block.

The procedure MakeLeaderBlock_i . When p_i wishes to form a new leader block b , it will run this procedure, by executing the following instructions:

1. Set $b.\text{type} := \text{lead}$, $b.\text{auth} := p_i$, $b.\text{view} := \text{view}_i$, $b.\text{slot} := \text{slot}_i(\text{lead})$.
2. Initially, set $b.\text{prev}$ to be the tips of Q_i .
3. Set $s := \text{slot}_i(\text{Tr})$ and $v := \text{view}_i$. If $s > 0$, then let $q \in Q_i$ be such that $q.\text{auth} = p_i$, $q.\text{type} = \text{lead}$ and $q.\text{slot} = s - 1$. If $b.\text{prev}$ does not already contain q , add q to this set.
4. If $h' = \max\{q.h : q \in b.\text{prev}\}$, then set $b.h := h' + 1$.
5. If p_i has not yet produced a block b with $b.\text{view} = \text{view}_i$ and $b.\text{type} = \text{lead}$ then:
 - a. Set $b.\text{just}$ to be a set of view v messages signed by $n - f$ processes in Π .
 - b. Set $b.1\text{-QC}$ to be a 1-QC in Q_i greater than or equal to all 1-QCs corresponding to messages in $b.\text{just}$.
6. If p_i has previously produced a block b with $b.\text{view} = \text{view}_i$ and $b.\text{type} = \text{lead}$ then let $q' \in Q_i$ be a 1-QC with $q'.\text{auth} = p_i$, $q'.\text{type} = \text{lead}$ and $q'.\text{slot} = s - 1$. Set $b.1\text{-QC} := q'$ and set $b.\text{just}$ to be the empty set.
7. Sign b with the values specified above, and send this block to all processes.
8. Set $\text{slot}_i(\text{lead}) := \text{slot}_i(\text{lead}) + 1$;

The pseudocode. The pseudocode appears in Algorithm 1 (with local variables described first, and the main code appearing later). Section 5.1 gives a ‘pseudocode walk-through’.

■ **Algorithm 1** Morpheus: local variables for p_i

```

1: Local variables
2:  $M_i$ , initially contains  $b_g$  and a 1-QC-certificate for  $b_g$  ▷ Automatically updated
3:  $Q_i$ , initially contains 1-QC-certificate for  $b_g$  ▷ Automatically updated
4:  $view_i$ , initially 0 ▷ The present view
5:  $slot_i(x)$  for  $x \in \{\text{lead}, \text{Tr}\}$ , initially 0 ▷ Present slot
6:  $voted_i(z, x, s, p_j)$  for  $z \in \{0, 1, 2\}$ ,  $x \in \{\text{lead}, \text{Tr}\}$ ,  $s \in \mathbb{N}_{\geq 0}$ ,  $p_j \in \Pi$ , initially 0
7:  $phase_i(v)$  for  $v \in \mathbb{N}_{\geq 0}$ , initially 0 ▷ The phase within the view
8: Other procedures and functions
9:  $lead(v)$  ▷ Leader of view  $v$ 
10:  $PayloadReady_i$  ▷ Set to 1 when ready to produce transaction block
11:  $MakeTrBlock_i$  ▷ Sends a new transaction block to all
12:  $LeaderReady_i$  ▷ Indicates whether ready to produce leader block
13:  $MakeLeaderBlock_i$  ▷ Sends a new leader block to all

```

5.1 Pseudocode walk-through

Lines 16-22: These lines are responsible for view changes. If p_i has received a quorum of end-view v messages for some greatest v greater than or equal to its present view, then it will use those to form a $(v + 1)$ -certificate and will send that certificate to all processes (immediately regarding that certificate as received and belonging to M_i). Upon seeing that it has received a v -certificate for some greatest view v greater than its present view, p_i will: (i) enter view v , (ii) send that v -certificate to all processes, and (iii) send a view v message to the leader of view v , along with any tips of Q_i corresponding to its own blocks. Process p_i will also do the same upon seeing q with $q.view$ greater than its present view: the latter action ensures that any block b produced by p_i during view v does not point to any b' with $b'.view > b.view$.

Lines 24-28. These lines are responsible for the production of 0-QCs. Upon producing any block, p_i sends it to all processes. Providing p_i is correct, meaning that the block is correctly formed etc, other processes will then send back a 0-vote for the block to p_i , who will form a 0-QC and send it to all processes.

Lines 30 and 31. These lines are responsible for producing new transaction blocks. Line 30 checks to see whether p_i is ready to produce a new transaction block, before line 31 produces the new block: $PayloadReady_i$ and $MakeTrBlock_i$ are specified in Section 5.

Lines 33 and 34. These lines are responsible for producing new leader blocks. Line 33 ensures that only the leader is asked to produce leader blocks, that it will only do so once ready (having received QCs for previous leader blocks, as required), and only when required to (only if Q_i does not have a single tip and if still in the first phase of the view). $LeaderReady_i$ and $MakeLeaderBlock_i$ are specified in Section 5.

Lines 36-47. These lines are responsible for determining when correct processes produce 1 and 2-votes for transaction blocks. Lines 36 and 37 dictate that no correct process produces 1 or 2-votes for transaction blocks while in view v until at least one leader block for the view has been finalized (according to the messages they have received), and only if there do not exist unfinalized leader blocks for the view. Given these conditions, p_i will produce a 1-vote for any transaction block b that is a single tip of M_i , so long as $b.1\text{-QC}$ is greater than or equal to any 1-QC it has seen. It will produce a 2-vote for a transaction block b if there exists q with $q.b = b$ which is a single tip of Q_i and if p_i has not seen any block of greater height. The latter condition is required to ensure that p_i cannot produce a 1-vote for some b' of greater height than b , and then produce a 2-vote for b (this fact is used in the proof of Theorem 2). After producing any 1 or 2-vote for a transaction block while in view v , p_i enters the second phase of the view and will no longer produce 1 or 2-votes for leader blocks while in view v .

■ **Algorithm 1** Morpheus: The instructions for p_i

14: Process p_i executes the following transitions at timeslot t (according to its local clock), until no further transitions apply. If multiple transitions apply simultaneously, then p_i executes the first that applies, before checking whether further transitions apply, and so on.

15: ▷ Update view

16: **If** there exists greatest $v \geq \text{view}_i$ s.t. M_i contains at least $f + 1$ end-view v messages **then**:

17: Form a $(v + 1)$ -certificate and send it to all processes;

18: **If** there exists some greatest $v > \text{view}_i$ such that either:

19: (i) M_i contains a v -certificate q , or (ii) Q_i contains q with $q.\text{view} = v$, **then**:

20: Set $\text{view}_i := v$; Send (either) q to all processes;

21: Send all tips q' of Q_i such that $q'.\text{auth} = p_i$ to $\text{lead}(v)$;

22: Send (v, q') signed by p_i to $\text{lead}(v)$, where q' is a maximal amongst 1-QCs seen by p_i

23: ▷ Send 0-votes and 0-QCs

24: **If** M_i contains some b s.t. $\text{voted}_i(0, b.\text{type}, b.\text{slot}, b.\text{auth}) = 0$:

25: Send a 0-vote for b (signed by p_i) to $b.\text{auth}$; Set $\text{voted}_i(0, b.\text{type}, b.\text{slot}, b.\text{auth}) := 1$;

26: **If** M_i contains a 0-quorum for some b s.t.:

27: (i) $b.\text{auth} = p_i$, **and** (ii) p_i has not previously sent a 0-QC for b to other processors, **then**:

28: Send a 0-QC for b to all processes;

29: ▷ Send out a new transaction block

30: **If** $\text{PayloadReady}_i = 1$ **then**:

31: MakeTrBlock_i ;

32: ▷ Send out a new leader block

33: **If** $p_i = \text{lead}(\text{view}_i)$, $\text{LeaderReady}_i = 1$, $\text{phase}_i(\text{view}_i) = 0$ **and** Q_i does not have a single tip:

34: MakeLeaderBlock_i ;

35: ▷ Send 1 and 2-votes for transaction blocks

36: **If** there exists $b \in M_i$ with $b.\text{type} = \text{lead}$ and $b.\text{view} = \text{view}_i$ **and**

37: there does not exist unfinalized $b \in M_i$ with $b.\text{type} = \text{lead}$ and $b.\text{view} = \text{view}_i$ **then**:

38: **If** there exists $b \in M_i$ with $b.\text{type} = \text{Tr}$, $b.\text{view} = \text{view}_i$ and which is a single tip of M_i s.t.:

39: (i) $b.1\text{-QC}$ is greater than or equal to every 1-QC in Q_i and;

40: (ii) $\text{voted}_i(1, \text{Tr}, b.\text{slot}, b.\text{auth}) = 0$, **then**:

41: Send a 1-vote for b to all processes; Set $\text{phase}_i(\text{view}_i) := 1$;

42: Set $\text{voted}_i(1, \text{Tr}, b.\text{slot}, b.\text{auth}) := 1$;

43: **If** there exists a 1-QC $q \in Q_i$ which is a single tip of Q_i s.t.:

44: (i) $q.\text{type} = \text{Tr}$ **and** (ii) $\text{voted}_i(2, \text{Tr}, q.\text{slot}, q.\text{auth}) = 0$, **then**:

45: **If** there does not exist $b \in M_i$ of height greater than $q.h$:

46: Send a 2-vote for $q.b$ to all processes; Set $\text{phase}_i(\text{view}_i) := 1$;

47: Set $\text{voted}_i(2, \text{Tr}, q.\text{slot}, q.\text{auth}) := 1$;

48: ▷ Vote for a leader block

49: **If** $\text{phase}(\text{view}_i) = 0$:

50: **If** $\exists b \in M_i$ with $b.\text{type} = \text{lead}$, $b.\text{view} = \text{view}_i$, $\text{voted}_i(1, \text{lead}, b.\text{slot}, b.\text{auth}) = 0$ **then**:

51: Send a 1-vote for b to all processes; Set $\text{voted}_i(1, \text{lead}, b.\text{slot}, b.\text{auth}) := 1$;

52: **If** $\exists q \in Q_i$ which is a 1-QC with $\text{voted}_i(2, \text{lead}, q.\text{slot}, q.\text{auth}) = 0$, $q.\text{type} = \text{lead}$,

53: $q.\text{view} = \text{view}_i$, **then**:

54: Send a 2-vote for $q.b$ to all processes; Set $\text{voted}_i(2, \text{lead}, q.\text{slot}, q.\text{auth}) := 1$;

55: ▷ Complain

56: **If** $\exists q \in Q_i$ which is maximal according to $\succeq$ amongst those that have not been finalized for time 6Δ since entering view view_i :

57: Send q to $\text{lead}(\text{view}_i)$ if not previously sent;

58: **If** $\exists q \in Q_i$ which has not been finalized for time 12Δ since entering view view_i :

59: Send the end-view message (view_i) signed by p_i to all processes;

519 **Lines 49–54.** These lines are responsible for determining when correct processes produce 1
 520 and 2-votes for leader blocks. Correct processes will only produce such votes while in the
 521 first phase of the view.

522 **Lines 56–59.** These lines are responsible for the production of new-view messages. The
 523 proof of Theorem 3 justifies the choice of 6Δ and 12Δ .

524 Proofs of consistency and liveness appear in Appendix A. A detailed analysis of latency
 525 and complexity appears in the online version of the paper, which can be found at <https://arxiv.org/abs/2502.08465>. Related work appears in Appendix B.

527 **6 Final Comments**

528 We have presented Morpheus Consensus, a protocol that dynamically adapts its struc-
 529 ture—shifting from a quiescent leaderless blockchain to an active leader-based DAG—while
 530 maintaining strong latency and complexity properties throughout. In high-throughput
 531 regimes, Morpheus demonstrates comparable performance to state-of-the-art DAG solutions
 532 like Autobahn. In low-throughput conditions, it achieves better latency and equivalent
 533 complexity versus established protocols such as PBFT and Tendermint, which fail to scale
 534 effectively under high load.

535 **References**

- 536 1 Nicolas Alhaddad, Sourav Das, Sisi Duan, Ling Ren, Mayank Varia, Zhuolun Xiang, and
 537 Haibin Zhang. Balanced byzantine reliable broadcast with near-optimal communication
 538 and improved computation. In *Proceedings of the 2022 ACM Symposium on Principles of*
 539 *Distributed Computing*, pages 399–417, 2022.
- 540 2 Balaji Arun, Zekun Li, Florian Suri-Payer, Sourav Das, and Alexander Spiegelman. Shoal++:
 541 High throughput dag bft can be fast! *arXiv preprint arXiv:2405.20488*, 2024.
- 542 3 Kushal Babel, Andrey Chursin, George Danezis, Anastasios Kichidis, Lefteris Kokoris-Kogias,
 543 Arun Koshy, Alberto Sonnino, and Mingwei Tian. Mysticeti: Reaching the limits of latency
 544 with uncertified dags. *arXiv preprint arXiv:2310.14821*, 2023.
- 545 4 Leemon Baird. The swirls hashgraph consensus algorithm: Fair, fast, byzantine fault tolerance.
 546 *Swirls Tech Reports SWIRLDS-TR-2016-01, Tech. Rep.*, 34:9–11, 2016.
- 547 5 Béla Bollobás. *Modern graph theory*, volume 184. Springer Science & Business Media, 2013.
- 548 6 Dan Boneh, Ben Lynn, and Hovav Shacham. Short signatures from the weil pairing. In
 549 *International conference on the theory and application of cryptology and information security*,
 550 pages 514–532. Springer, 2001.
- 551 7 Gabriel Bracha. Asynchronous byzantine agreement protocols. *Information and Computation*,
 552 75(2):130–143, 1987.
- 553 8 Ethan Buchman. *Tendermint: Byzantine fault tolerance in the age of blockchains*. PhD thesis,
 554 2016.
- 555 9 Ethan Buchman, Jae Kwon, and Zarko Milosevic. The latest gossip on bft consensus. *arXiv*
 556 *preprint arXiv:1807.04938*, 2018.
- 557 10 Miguel Castro, Barbara Liskov, et al. Practical byzantine fault tolerance. In *OSDI*, volume 99,
 558 pages 173–186, 1999.
- 559 11 Xiaohai Dai, Guanxiong Wang, Jiang Xiao, Zhengxuan Guo, Rui Hao, Xia Xie, and Hai Jin.
 560 Lightdag: A low-latency dag-based bft consensus through lightweight broadcast. *Cryptology*
 561 *ePrint Archive*, 2024.
- 562 12 Xiaohai Dai, Zhaonan Zhang, Jiang Xiao, Jingtao Yue, Xia Xie, and Hai Jin. Gradeddag:
 563 An asynchronous dag-based bft consensus with lower latency. In *2023 42nd International*
 564 *Symposium on Reliable Distributed Systems (SRDS)*, pages 107–117. IEEE, 2023.

- 565 13 George Danezis, Lefteris Kokoris-Kogias, Alberto Sonnino, and Alexander Spiegelman. Narwhal
566 and tusk: a dag-based mempool and efficient bft consensus. In *Proceedings of the Seventeenth*
567 *European Conference on Computer Systems*, pages 34–50, 2022.
- 568 14 Adam Gagol, Damian Leśniak, Damian Straszak, and Michał Świątek. Aleph: Efficient atomic
569 broadcast in asynchronous networks with byzantine nodes. In *Proceedings of the 1st ACM*
570 *Conference on Advances in Financial Technologies*, pages 214–228, 2019.
- 571 15 Neil Girdharan, Florian Suri-Payer, Ittai Abraham, Lorenzo Alvisi, and Natacha Crooks.
572 Autobahn: Seamless high speed bft. In *Proceedings of the ACM SIGOPS 30th Symposium on*
573 *Operating Systems Principles*, pages 1–23, 2024.
- 574 16 Guy Golan Gueta, Ittai Abraham, Shelly Grossman, Dahlia Malkhi, Benny Pinkas, Michael
575 Reiter, Dragos-Adrian Seredinschi, Orr Tamir, and Alin Tomescu. Sbft: A scalable and
576 decentralized trust infrastructure. In *2019 49th Annual IEEE/IFIP international conference*
577 *on dependable systems and networks (DSN)*, pages 568–580. IEEE, 2019.
- 578 17 Idit Keidar, Eleftherios Kokoris-Kogias, Oded Naor, and Alexander Spiegelman. All you need
579 is dag. In *Proceedings of the 2021 ACM Symposium on Principles of Distributed Computing*,
580 pages 165–175, 2021.
- 581 18 Idit Keidar, Andrew Lewis-Pye, and Ehud Shapiro. Grassroots consensus. *arXiv preprint*
582 *arXiv:2505.19216*, 2025.
- 583 19 Idit Keidar, Oded Naor, Ouri Poupko, and Ehud Shapiro. Cordial miners: Fast and efficient
584 consensus for every eventuality. *arXiv preprint arXiv:2205.09174*, 2022.
- 585 20 Ramakrishna Kotla, Lorenzo Alvisi, Mike Dahlin, Allen Clement, and Edmund Wong. Zyzzyva:
586 speculative byzantine fault tolerance. In *Proceedings of twenty-first ACM SIGOPS symposium*
587 *on Operating systems principles*, pages 45–58, 2007.
- 588 21 Andrew Miller, Yu Xia, Kyle Croman, Elaine Shi, and Dawn Song. The honey badger
589 of bft protocols. In *Proceedings of the 2016 ACM SIGSAC conference on computer and*
590 *communications security*, pages 31–42, 2016.
- 591 22 Kartik Nayak, Ling Ren, Elaine Shi, Nitin H Vaidya, and Zhuolun Xiang. Improved extension
592 protocols for byzantine broadcast and agreement. *arXiv preprint arXiv:2002.11321*, 2020.
- 593 23 Victor Shoup. Practical threshold signatures. In *International Conference on the Theory and*
594 *Applications of Cryptographic Techniques*, pages 207–220. Springer, 2000.
- 595 24 Nibesh Shrestha, Rohan Shrothrium, Aniket Kate, and Kartik Nayak. Sailfish: Towards
596 improving latency of dag-based bft. *Cryptology ePrint Archive*, 2024.
- 597 25 Yonatan Sompolinsky, Yoad Lewenberg, and Aviv Zohar. Spectre: A fast and scalable
598 cryptocurrency protocol. *Cryptology ePrint Archive*, 2016.
- 599 26 Alexander Spiegelman, Balaji Arun, Rati Gelashvili, and Zekun Li. Shoal: Improving dag-bft
600 latency and robustness. *arXiv preprint arXiv:2306.03058*, 2023.
- 601 27 Alexander Spiegelman, Neil Girdharan, Alberto Sonnino, and Lefteris Kokoris-Kogias. Bull-
602 shark: Dag bft protocols made practical. In *Proceedings of the 2022 ACM SIGSAC Conference*
603 *on Computer and Communications Security*, pages 2705–2718, 2022.
- 604 28 TK Srikanth and Sam Toueg. Simulating authenticated broadcasts to derive simple fault-
605 tolerant algorithms. *Distributed Computing*, 2(2):80–94, 1987.
- 606 29 Maofan Yin, Dahlia Malkhi, Michael K Reiter, Guy Golan Gueta, and Ittai Abraham. Hotstuff:
607 Bft consensus in the lens of blockchain. *arXiv preprint arXiv:1803.05069*, 2018.
- 608 30 Maofan Yin, Dahlia Malkhi, Michael K Reiter, Guy Golan Gueta, and Ittai Abraham. Hotstuff:
609 Bft consensus with linearity and responsiveness. In *Proceedings of the 2019 ACM Symposium*
610 *on Principles of Distributed Computing*, pages 347–356, 2019.

611 **A** Establishing consistency and liveness

612 Let M^* be the set of all messages received by any process during the execution. Towards
613 establishing consistency, we first prove the following lemma.

614 ► **Lemma 1.** *If $q, q' \in M^*$ are 1-QCs with $q \leq q'$ and $q' \leq q$, then $q.b = q'.b$.*

615 **Proof.** Suppose $q.view = q'.view$, $q.type = q'.type$, and $q.h = q'.h$. Consider first the
 616 case that $q.b$ and $q'.b$ are both leader blocks for the same view. If $q.slot = q'.slot$, but
 617 $q.b \neq q'.b$, then no correct process can produce 1-votes for both blocks. This gives an
 618 immediate contradiction, since two subsets of Π of size $n - f$ must have a correct process
 619 in the intersection, meaning that 1-QCs cannot be produced for both blocks. So, suppose
 620 that $q'.slot > q.slot$. Since each leader block b with $b.slot = s > 0$ must point to a leader
 621 block b' with $b'.auth = b.auth$ and $b'.slot = s - 1$, it follows that $q'.h > q.h$, which also gives
 622 a contradiction.

623 So, consider next the case that $q.b$ and $q'.b$ are distinct transaction blocks. Since both
 624 blocks are of the same height, and since any correct process only votes for a block when it is
 625 a sole tip of its local value M_i , no correct process can vote for both blocks. Once again, this
 626 gives the required contradiction. ◀

627 Note that Lemma 1 also suffices to establish a similar result for 2-QCs, since no block
 628 can receive a 2-QC without first receiving a 1-QC: No correct process produces a 2-vote for
 629 any block without first receiving a 1-QC for the block.

630 Lemma 1 suffices to show that we can think of all 1-QCs $q \in M^*$ as belonging to a
 631 hierarchy, ordered by $q.view$, then by $q.type$, and then by $q.h$, such that if q and q' belong to
 632 the same level of this hierarchy then $q.b = q'.b$.

633 ► **Theorem 2.** *The Morpheus protocol satisfies consistency.*

634 **Proof.** Given the definition of $\mathcal{F}$ from Section 4, let us say $b' \rightarrow b$ iff:

- 635 ■ $b' = b$, or;
- 636 ■ $b' \neq b_g$ and $b'' \rightarrow b$, where $q = b'.1\text{-QC}$ and $b'' = q.b$.

637 To establish consistency it suffices to show the following:

638 (†): If b has a 1-QC $q_1 \in M^*$ and also a 2-QC $q_2 \in M^*$, then for any 1-QC $q \in M^*$ such that
 639 $q \geq q_1$, $q.b \rightarrow b$.

640 Given (†), suppose $M_1 \subseteq M_2 \subseteq M^*$. For each $i \in \{1, 2\}$, let M'_i be the largest set of
 641 blocks in M_i that is downward closed (in the sense specified in Section 4). Let q'_i be a
 642 maximal 2-QC in M_i such that $q'_i.b \in M'_i$, and set $b_i^* = q'_i.b$, or if there is no such 2-QC in
 643 M_i , set $b_i^* = b_g$. Let the sequence $b_k, \dots, b_1 = b_g$ be such that $b_k = b_2^*$, and, for each $j < k$, if
 644 $q = b_{j+1}.1\text{-QC}$, then $q.b = b_j$. From (†) it follows that b_1^* belongs to the sequence $b_k, \dots, b_1$,
 645 so that $\mathcal{F}(M_2) \supseteq \mathcal{F}(M_1)$.

646 We establish (†) by induction on the level of the hierarchy to which q belongs. If $q \leq q_1$
 647 (and $q_1 \leq q$) then the result follows from Lemma 1.

648 For the induction step, suppose that $q > q_1$ and suppose first that $q.type = \text{lead}$.
 649 Let $s = q.slot$, $v = q.view$. By validity of $q.b$, if $s > 0$, $q.b$ points to a unique b^* with
 650 $b^*.type = \text{lead}$, $b^*.auth = q.auth$ and $b^*.slot = s - 1$. If $s = 0$ or $b^*.view < v$, then $q.just$ (i.e.
 651 $(q.b).just$) contains $n - f$ view v messages, each signed by a different process in Π . Note that,
 652 in this case, any correct process that produces a 2-vote for b must do so before sending a view
 653 v message. It follows that, in this case, $q.1\text{-QC}$ (i.e. $(q.b).1\text{-QC}$) belongs to a level of the
 654 hierarchy strictly below q and greater than or equal to that of q_1 . The result therefore follows
 655 by the induction hypothesis. If $s > 0$ and $b^*.view = v$, then $q.1\text{-QC}$ is a 1-QC-certificate for

656 b^* . Once again, q .1-QC therefore belongs to a level of the hierarchy strictly below q and
 657 greater than or equal to that of q_1 , so that the result follows by the induction hypothesis.

658 So, suppose next that $q.type = Tr$. Note that, in this case, any correct process that
 659 produces a 2-vote for b must do so before sending a 1-vote for $q.b$. If $q.view > b.view$ this
 660 follows immediately, because a correct process p_i only sends 1 or 2-votes for any block b'
 661 while $view_i = b'.view$. If $q.view = b.view$ and $b.type = lead$, this follows because no correct
 662 process sends 1 or 2-votes for a leader block after having voted for a transaction block within
 663 the same view. If $q.view = b.view$ and $b.type = Tr$, this follows because any correct process
 664 only sends a 2-vote for b so long as there does not exist $b' \in M_i$ of height greater than b .
 665 Also, any correct process that produces a 2-vote for b will not vote for $q.b$ unless q .1-QC is
 666 greater than or equal to any 1-QC it has received. It follows that q .1-QC belongs to a level
 667 of the hierarchy strictly below q and greater than or equal to that of q_1 . Once again, the
 668 result follows by the induction hypothesis. ◀

669 ▶ **Theorem 3.** *The Morpheus protocol satisfies liveness.*

670 **Proof.** Towards a contradiction, suppose that correct p_i produces a transaction block b ,
 671 which never becomes finalized (according to the messages received by p_i). Note that all
 672 correct processes eventually send 0-votes for b to p_i , meaning that p_i forms a 0-certificate for
 673 b , which is eventually received by all correct processes. Since correct processes send end-view
 674 messages if a some QC is not finalized for sufficiently long within any given view (see line
 675 58), correct processes must therefore enter infinitely many views. Let v be a view with a
 676 correct leader, such that the first correct process p_j to enter view v does so at some timeslot
 677 after GST, and after p_i produces b . Process p_j sends a v-certificate to all processes upon
 678 entering the view, meaning that all correct processes enter the view within time Δ of p_j
 679 doing so. Upon entering view v , at time t say, note that p_i will send a QC for a transaction
 680 block b' that it has produced to the leader. This block b' has a slot number greater than or
 681 equal to that of b . The leader will produce a leader block observing b' by time $t + 3\Delta$, which
 682 will be finalized (according to the messages received by p_i) by time $t + 6\Delta$. ◀

683 B Related Work

684 Morpheus uses a PBFT [10] style approach to view changes, while consistency between
 685 finalised transaction blocks within the same view uses an approach similar to Tendermint
 686 [8, 9] and Hotstuff [30]. Hotstuff's approach of relaying all messages via the leader could
 687 be used by Morpheus during low throughput to decrease communication complexity, but
 688 this is unlikely to lead to a decrease in 'real' latency (i.e. actual finalisation times). The
 689 optimistic 'fast commit' of Zyzzyva [16, 20] can also be applied as a further optimisation. The
 690 recent paper [18] shows how to implement player reconfiguration for a form of the Morpheus
 691 protocol.

692 Morpheus transitions between being a leaderless 'linear' blockchain during low throughput
 693 to a leader-based DAG-protocol during high throughput. DAG protocols have been studied for
 694 a number of years, Hashgraph [4] being an early example. Hashgraph builds an unstructured
 695 DAG and suffers from latency exponential in the number of processes. Spectre was another
 696 early DAG protocol, designed for the 'permissionless' setting [25], with proof-of-work as
 697 the mechanism for sybil resistance. The protocol implements a 'payment system', but does
 698 not totally order transactions. Aleph [14] is more similar to most recent DAG protocols in
 699 that it builds a structured DAG in which each process proceeds to the next 'round' after

700 receiving blocks from $2f + 1$ processes corresponding to the previous round, but still has
701 greater latency than modern DAG protocols.

702 More recent DAG protocols use a variety of approaches to consensus. Narwhal [13] builds
703 a DAG for the purpose of ensuring data availability, from which (one option is that) a
704 protocol like Hotstuff or PBFT can then be used to efficiently establish a total ordering on
705 transactions. DAG-Rider [17], on the other hand, builds the DAG in such a way that a total
706 ordering can be extracted from the structure of the DAG, with zero further communication
707 cost. The protocol proceeds in ‘waves’, where each wave consists of four rounds, each round
708 building one ‘layer’ of the DAG. In each round, each process uses an instance of Reliable
709 Broadcast (RBC) to disseminate their block for the round. Each wave has a leader and an
710 expected six rounds (6 sequential RBCs) are required to finalise the leader’s block for the
711 first round of the wave. This finalises all blocks observed by that leader block, but other
712 blocks (such as those in the same round as the leader block) may have significantly greater
713 latency. Tusk [13] is an implementation based on DAG-Rider.

714 Given the ability of DAG-Rider to handle significantly higher throughput in many settings,
715 when compared to protocols like PBFT that build a linear blockchain, much subsequent
716 work has taken a similar approach, while looking to improve on latency. While DAG-Rider
717 functions in asynchrony, Bullshark [27] is designed to achieve lower latency in the partially
718 synchronous setting. GradedDAG [12] and LightDAG [11] function in asynchrony, but look to
719 improve latency by replacing RBC [7] with weaker primitives, such as consistent broadcast [28].
720 This means that those protocols solve Extractable SMR (as defined in Section 3), rather than
721 SMR, and that further communication may be required to ensure full block dissemination in
722 executions with faulty processes. Cordial Miners [19] has versions for both partial synchrony
723 and asynchrony and further decreases latency by using the DAG structure (rather than any
724 primitive such as Consistent or Reliable Broadcast) for equivocation exclusion. Mysticeti
725 [3] builds on Cordial Miners and establishes a mechanism to accommodate multiple leaders
726 within a single round. Shoal [26] and Shoal++ [2] extend Bullshark by establishing a
727 ‘pipelining approach’ that implements simultaneous instances of Bullshark with a leader in
728 each round. This reduces latency in the good case because one is required to wait less time
729 before reaching a round in which a leader block is finalised. Both of these papers, however,
730 use a ‘reputation’ system to select leaders, which comes with its own trade-offs. Sailfish [24]
731 similarly describes a mechanism where each round has a leader, but does not make use of a
732 reputation system. As noted previously, the protocol most similar to Morpheus during high
733 throughput is Autobahn [15]. One of the major distinctions between Autobahn and those
734 previously discussed, is that most blocks are only required to point to a single parent. This
735 significantly decreases communication complexity when the number of processes is large and
736 allows one to achieve linear amortised communication complexity without the use of erasure
737 coding [1, 22] or batching [21].