

GoToCloud Optimization of Cloud Computing Environment for Accelerating Cryo-EM Structure-Based Drug Design.

Corresponding Author: Dr Toshio Moriya

Version 0:

Reviewer comments:

Reviewer #1

(Remarks to the Author)

The authors in this manuscript present a software package written for structural biology end users to leverage public cloud resources. The authors nicely demonstrate the tradeoff between cost and speed, arriving at specific recommendations for instance type for specific RELION jobs.

The paper is clearly written, I don't have any suggestions for the text or figures.

The only minor comment relates to the code in the online Github repository. Can the authors include a README file with tutorial information on how to use their software? The more information they can provide, the better for new users. Additionally, can the authors provide a license file in the Github repo to clearly indicate under what license the software is released?

Reviewer #2

(Remarks to the Author)

Brief Summary:

The manuscript describes a way to configure Amazon Web Services (AWS) to run Cryo-EM single particle analysis on multiple nodes. The authors use shell scripts to set up the processing environment with AWS products, like Elastic Compute Cloud and ParallelCluster for managing the computing cluster, account management tools for security and billing, and S3 and Elastic File System for storing data. Relion, the data processing software, is compiled and tuned for the hardware on the GPU and CPU instances. This is a way to process the large amounts of data (micrograph movies) generated by experiments and in the future could be a way to efficiently run structure based drug design workflows on AWS. The benefit of using AWS or any other cloud computing provider is that users do not need to maintain their own hardware and can dynamically adjust the computational resources depending on the need. Benchmarks for different processing stages are presented and guidance on how to select computing resources based on processing time and cost is provided.

Overall Impression:

As data sizes continue to grow in Cryo-EM and other experimental methods, computational resources and the expertise in managing these resources and associated data processing software need to be more easily accessible to the people performing the experiments. While there are other approaches that also use cloud computing (as cited by the authors), the approach presented in this manuscript supports multiple nodes for a single processing job. The other cited approaches only support single node (but still multi-GPU) or the other approach that also supports multiple nodes has not been updated for several years ("cryoem-cloud-tools", <https://github.com/cianfrocco-lab/cryoem-cloud-tools>). These wrappers for simplifying the deployment of software and hardware resources lowers the barrier for running these computationally expensive jobs.

Moreover, with the way accounts are set up with this approach, it would be possible for a microscope facility to maintain the data processing software stack and have it available for users who have separate accounts for billing their use of the computational resources. With the granularity of AWS billing, users would really only pay for their usage of AWS resources. And since the authors collaborated with AWS staff, it would be much easier for the microscope facility to maintain and update these scripts to work with updates or improvements in AWS. This approach can centralize the technical aspects of maintaining and optimizing the data processing software, and with optimizing the use of AWS resources, letting the users focus on processing their data.

Specific Comments:

- 1) Why were the g5 and g4dn instances chosen? Why not p4 (A100) or p5 (H100)? The p series is for general GPU compute. g5 and g4dn seem to be more for machine learning and graphics intensive workflows.
- 2) Why was the GPU Relion executable built with CUDA_ARCH=75 (from the cmake flags)? For the NVIDIA A10G, it can be 86. Can any benchmarking be done comparing the performance of the A10G GPU with just the CUDA_ARCH=75 optimization and the performance with the CUDA_ARCH=86 optimization? Also, it should be possible to build multiple architecture support into the executable (nvcc does support this, I'm not sure how to modify cmake to do this). Or is this flag separate from optimization? The architecture flag should include some architecture-specific improvements.
- 3) For completeness, and this can be done in the supplementary materials,
 - 3a) Can you add a sentence to describe the linux distribution used? From the scripts, it looks like Ubuntu 20.04.
 - 3b) What were the sizes (disk space) of the benchmark datasets? That can be added to Table S1. This would give readers an idea of the dataset sizes.
- 4) Did the reported costs include storage costs (S3, EFS, Lustre)? From the spreadsheet, it looks like it is the processing costs only. Can you provide a breakdown of the total costs in terms of processing time (compute) and storage?
- 5) There are hardcoded S3 bucket URLs in scripts (e.g. https://kek-gtc-master-s3-bucket.s3.ap-northeast-1.amazonaws.com/gtc_efs_setting.json in `gtc_setup_gotocloud_environment.sh`). Is it intended for the scripts to have hardcoded URLs for your S3 bucket? Are the scripts useable by other people to set up their own clusters? Also, the S3 URLs to be publicly accessible?
- 6) How long does it take to transfer data to the S3 bucket? Are you transferring from Japan to an AWS region in the United States?
- 7) Also, how long does it take to transfer data from the S3 bucket to the Amazon FSx Lustre filesystem?

Reviewer #3

(Remarks to the Author)

The manuscript concerns the computational requirements of Single Particle Analysis for cryoEM. These can be high, and it is certainly true that structural biologists often use inappropriate setups and hardware, and may do more processing than necessary. Target resolution is also a key quantity, and achieving the highest resolution possible is not always necessary for the biological question.

Many facilities and groups have sought to optimise the use of compute resources. Some efforts have been published, most not. Relion has been benchmarked several times before. Cianfrocco (ref 16) highlighted the benefits of spot instances on AWS in 2015. The question is whether this article provides any additional insights or a practical way for others to manage their resources.

GoToCloud supports multi-node processing using a selection of common cryoEM software via AWS pcluster. It manages the hardware selection, dependencies, and cryoEM software installation. It also handles configuration with respect to AWS Regions. It would appear to be a good solution for the authors, but it is unclear to me to what extent the code or the article are useful to others.

Pre-installation of a range of executables is handled by a management account. It was unclear at first whether the authors were providing a service to the community, but I believe they are only providing the tools to do this. It is not even clear if it is being used at the Photon Factory. Maintaining up-to-date software is a headache for all facility and compute managers. Software gets rapidly out-of-date, as shown by their usage of Relion 3 and 4 (version 5 is now out).

Their benchmarks show that the cost/performance can be optimised for each job type. They choose a criterion of closest to the origin of a cost vs time graph, and I think that is reasonable. They show that this is dependent on instance type, and will need to be updated for as new hardware comes along. More urgently, I would assume that it is dependent on the dataset. They have only benchmarked this calculation on a single example (nitrite reductase) and so the variation in their conclusions is not tested. They do not describe what would happen in normal usage. Does this benchmarking have to be done on each new dataset in order to determine the best configuration? Is the user expected to do this? Can anything be inferred from short runs (probably not as Relion searches tend to be adaptive)?

The test against resolution was not what I was expecting. They record the computational requirements as the structure solution progresses through the different stages, rather than testing against adjustable parameters (such as number of particles picked, search parameters, etc). Possibly the only practical decision to be made is how many cycles of particle polishing to do. Here, the conclusion is a bit strange. After emphasizing the substantial cost of these additional cycles, they conclude it is worth doing anyway. They suggest that the modest resolution gain will help model building, although this is not actually tried. Overall, this test seems a bit pointless, unless they are going to demonstrate some kind of stopping condition. Again, this test is done on a single benchmark, so it is not clear how general it is.

The core of GoToCloud seems to be a set of GTC configuration scripts. These are available on github, and are a set of bash scripts. I could see no documentation there, and I think it would be very difficult for anyone else to set up GoToCloud, let alone adapt it or update it. There is some explanation in the Supplementary Methods which helps to give the idea, but not enough for a 3rd party to adopt this solution.

Some minor comments:

1. Figures 2 to 5 show the same plots for 4 different job types. I don't think this is necessary and some space could be saved.
2. Figure 1a contains a typo - it should be UCSF Chimera(X).

In conclusion, although the work is technically correct, I don't think there is much novelty in this article. The benchmarks and conclusions are specific to the two examples studied, and there have been other benchmarks published. If it is meant to be a practical solution, then there are many questions around how someone would deploy the GoToCloud system, and how users would use it for their datasets. I cannot recommend publication.

Author Rebuttal letter:

1

2 Response to Reviewers

3

4 Reviewers' comments:

5

6 Responses to Reviewer #1

7 C1-1) The authors in this manuscript present a software package written for structural
8 biology end users to leverage public cloud resources. The authors nicely demonstrate the
9 tradeoff between cost and speed, arriving at specific recommendations for instance type
10 for specific RELION jobs.

11 The paper is clearly written, I don't have any suggestions for the text or figures.

12

13 The only minor comment relates to the code in the online Github repository. Can the
14 authors include a README file with tutorial information on how to use their software?

15 The more information they can provide, the better for new users. Additionally, can the
16 authors provide a license file in the Github repo to clearly indicate under what license the
17 software is released?

18

19 Response 1-1) Thank you for your encouraging feedback on our manuscript and for
20 taking the time to review the associated GitHub repository. We appreciate your valuable
21 suggestion regarding the inclusion of a README file and a license file. We agree that
22 providing additional information for new users is crucial to significantly enhance the
23 usability and accessibility of our code for the research community. Therefore, we have
24 included both a README file and a license file in the GitHub repository

25 (<https://github.com/KEK-SBRC-CryoEM/gotocloud>) and added the following sentences
26 in line 529-534 of the main text: "The README file in the GitHub repository provides
27 step-by-step instructions on how to use the GoToCloud scripts for setting up the
28 GoToCloud platform and creating the instances on user's AWS account. It also includes
29 links to websites where readers can find detailed versions of these procedures,
30 comprehensive tutorials on using the GoToCloud scripts, and installation instructions for
31 the main software supported by the GoToCloud platform." The license file is for the GNU
32 GENERAL PUBLIC LICENSE Version 3, which clearly outlines the terms under which
33 our GoToCloud scripts can be used, modified, and distributed. Thank you once again for
34 your feedback and for helping us improve the accessibility and usability of our software.

35

1

36 Responses to Reviewer #2

37

38 Brief Summary:

39 The manuscript describes a way to configure Amazon Web Services (AWS) to run Cryo-
40 EM single particle analysis on multiple nodes. The authors use shell scripts to set up the
41 processing environment with AWS products, like Elastic Compute Cloud and
42 ParallelCluster for managing the computing cluster, account management tools for
43 security and billing, and S3 and Elastic File System for storing data. Relion, the data
44 processing software, is compiled and tuned for the hardware on the GPU and CPU
45 instances This is a way to process the large amounts of data (micrograph movies)
46 generated by experiments and in the future could be a way to efficiently run structure
47 based drug design workflows on AWS. The benefit of using AWS or any other cloud
48 computing provider is that users do not need to maintain their own hardware and can
49 dynamically adjust the computational resources depending on the need. Benchmarks for
50 different processing stages are presented and guidance on how the select computing
51 resources based on processing time and cost is provided.

52

53 Overall Impression:

54 As data sizes continue to grow in Cryo-EM and other experimental methods,

55 computational resources and the expertise in managing these resources and associated
56 data processing software need to be more easily accessible to the people performing the
57 experiments. While there are other approaches that also use cloud computing (as cited by
58 the authors), the approach presented in this manuscript supports multiple nodes for a
59 single processing job. The other cited approaches only support single node (but still multi-
60 GPU) or the other approach that also supports multiple nodes has not been updated for
61 several years ("cryoem-cloud-tools", [https://github.com/cianfrocco-lab/cryoem-cloud-](https://github.com/cianfrocco-lab/cryoem-cloud-tools)
62 [tools](https://github.com/cianfrocco-lab/cryoem-cloud-tools)). These wrappers for simplifying the deployment of software and hardware
63 resources lowers the barrier for running these computationally expensive jobs.

64
65 Moreover, with the way accounts are set up with this approach, it would be possible for
66 a microscope facility to maintain the data processing software stack and have it available
67 for users who have separate accounts for billing their use of the computational resources.
68 With the granularity of AWS billing, users would really only pay for their usage of AWS
69 resources. And since the authors collaborated with AWS staff, it would be much easier
70 for the microscope facility to maintain and update these scripts to work with updates or
71 improvements in AWS. This approach can centralize the technical aspects of maintaining

2
72 and optimizing the data processing software, and with optimizing the use of AWS
73 resources, letting the users focus on processing their data.

74
75 Thank you for your accurate and concise summary of our current work. Your expert
76 comments on AWS were particularly valuable in enhancing the quality of our current
77 paper.

78
79 Specific Comments:

80 C2-1) Why were the g5 and g4dn instances chosen? Why not p4 (A100) or p5 (H100)?
81 The p series is for general GPU compute. g5 and g4dn seem to be more for machine
82 learning and graphics intensive workflows.

83
84 Response 2-1) Thank you for pointing out that we have omitted an explanation of why
85 we chose G4dn and G5 over P4 and P5. We added the following sentences in line 198-
86 207 of the main text: "The P4 (p4d.24xlarge) and P5 (p5.48xlarge) instances were also
87 considered, but we chose the G4dn and G5 instance types over P4 and P5 primarily
88 because of their flexibilities in GPU card numbers and hourly rates that allow us to
89 optimize number of GPUs (i.e. less than 8 GPU cards) and so reduce the cost for each
90 RELION job type. While P4 and P5 instance types support only the 8 GPU configuration,
91 G4dn and G5 instance types offer configurations with 1, 4, and 8 GPU cards. Additionally,
92 the hourly rates for P4 (32.7726 USD) and P5 (98.32 USD) instances are significantly
93 higher than those for G4dn (USD 7.824 for g4dn.metal) and G5 (USD 16.288 for
94 g5.48xlarge) instances with 8 GPU cards. Therefore, finer adjustments for cost
95 performance are achievable with G4dn and G5."

96
97 C2-2) Why was the GPU Relion executable built with CUDA_ARCH=75 (from the
98 cmake flags)? For the NVIDIA A10G, it can be 86. Can any benchmarking be done
99 comparing the performance of the A10G GPU with just the CUDA_ARCH=75
100 optimization and the performance with the CUDA_ARCH=86 optimization? Also, it
101 should be possible to build multiple architecture support into the executable (nvcc does
102 support this, I'm not sure how to modify cmake to do this). Or is this flag separate from
103 optimization? The architecture flag should include some architecture- specific
104 improvements.

105
106 Response 2-2) We are grateful for your suggestion of additional optimization possibilities
107 for the G5 instance types equipped with NVIDIA A10G Tensor Core GPUs. For this, we

3
108 additionally built a RELION4.0 executable with CUDA_ARCH=86 and conducted
109 benchmarks for Refine3D, Class3D, and Class2D using the G5 instances. The results
110 showed clear improvements in execution times and costs for Refine3D, some
111 improvement for Class3D, and no improvement for Class2D. The improvements become
112 more pronoun as the number of nodes increases. Accordingly, we added the "Results of
113 a further optimization for the G5 instance types" subsection in line 306-315 of the main
114 text as follow:

115 "Results of a further optimization for the G5 instance types:
116 A further optimization for the G5 instance types equipped with NVIDIA A10G Tensor
117 Core GPUs were tempted. We built another RELION excitable using the cmake

118 configuration optimal for the compute capability of the CUDA architecture of A10G (i.e.,
119 “-DCUDA_ARCH=86”), in addition to the one which is optimal for the G4dn with
120 NVIDIA T4 Tensor Core GPU (i.e., “-DCUDA_ARCH=75”), and conducted
121 benchmarks for Refine3D, Class3D, and Class2D using the G5 instance types. The results
122 showed clear improvements in execution times and costs for Refine3D, slight
123 improvement for Class3D, and no improvement for Class2D (Fig. S6). The
124 improvements become more pronoun as the number of nodes increases.”
125 Also, figure S6 has been included in the Supplemental Information to present these
126 results, and raw data has been appended to the Refine3D, Class3D, and Class2D sheets
127 in Supplemental Data 2. To provide the details on compilation settings and the types of
128 RELION jobs applied, we added the “EXE04_GPU” entry in Supplemental Data 1.
129
130 C2-3) For completeness, and this can be done in the supplementary materials,
131 C2-3a) Can you add a sentence to describe the linux distribution used? From the scripts,
132 it looks like Ubuntu 20.04.
133
134 Response 2-3a) Thank you for pointing out that the description of the Linux distribution
135 was missing. We added the following sentence in line 174-175 of the Supplemental
136 Information: “Currently, the Linux distribution of this plucster instance is set to Ubuntu
137 20.04 LTS (Long Term Support) in the pcluster configuration file.”
138
139 C2-3b) What were the sizes (disk space) of the benchmark datasets? That can be added
140 to Table S1. This would give readers an idea of the dataset sizes.
141

4

142 Response 2-3b) We appreciate your suggestion. We have added the sizes of the
143 benchmark datasets (EMPIAR-10581: 1.0 TB (MRC), EMPIAR-10641: 373.1 GB
144 (TIFF)) to Table S1 as requested.
145
146 C2-4) Did the reported costs include storage costs (S3, EFS, Lustre)? From the
147 spreadsheet, it looks like it is the processing costs only. Can you provide a breakdown of
148 the total costs in terms of processing time (compute) and storage?
149
150 Response 2-4) We appreciate you bringing the significance of AWS storage costs to our
151 attention. We agree that this topic is of great interest to readers. As you pointed out, the
152 reported figures were the processing costs only. To address this point, we added a
153 breakdown of all the associated costs and calculated the “Grand Total Cost”, which
154 includes the storage costs, in Supplemental Data 2. The columns for the costs of the
155 pcluster head node, Lustre, S3 and EFS associated to each benchmark run are added.
156 Since the pricing of the S3 and EFS provided by AWS are listed as a monthly rate per
157 GB (USD/GB/month), we first checked the billing statement of our AWS account for the
158 month when we conducted the benchmark tests (February 2023), and then extracted the
159 costs associated to the S3 and EFS. Newly added the “Pricing” sheet in Supplemental
160 Data 2 provides these costs. For the cost calculation of each RELION job execution, these
161 costs in the month are further multiple by the ratio of the process time (in hours) relative
162 to the period of the month. For this, the following paragraph were added to the main text
163 (line 217-224): “As reference for readers, the Supplemental Data 2 also contains a
164 breakdown of all the associated costs as well as a “grand total cost”, including the storage
165 costs. Since the pricing of the S3 and EFS provided by AWS are listed as a monthly rate
166 per GB (USD/GB/month), the costs associated to the S3 and EFS were extracted from the
167 billing statement of our AWS account for the month (February 2023) when the
168 benchmark tests were conducted (provide in the “Pricing” sheet). For the calculation of
169 the ground total cost for each RELION job execution, these storage costs in the month
170 are further multiple by the ratio of the process time (in hours) relative to a month.”
171
172 C2-5) There are hardcoded S3 bucket URLs in scripts (e.g. [https://kek-gtc-master-s3-](https://kek-gtc-master-s3-bucket.s3.ap-northeast-1.amazonaws.com/gtc_efs_setting.json)
173 [bucket.s3.ap-northeast-1. amazonaws.com/gtc_efs_setting.json](https://kek-gtc-master-s3-bucket.s3.ap-northeast-1.amazonaws.com/gtc_efs_setting.json) in
174 [gtc_setup_gotocloud_environment.sh](#)). Is it intended for the scripts to have hardcoded
175 URLs for your S3 bucket? Are the scripts useable by other people to set up their own
176 clusters? Also, the S3 URLs to be publicly accessible?
177

5

178 Response 2-5) Thank you for taking the time to review the GoToCloud scripts and

179 finding the problem. The command interface of `gtc_setup_gotocloud_environment.sh` is
180 modified in the GoToCloud GitHub. Now, the command has an option “-o
181 OWN_SHARED_S3_BUCKET”, with which user can specify the path to a shared S3
182 bucket of user’s own GoToCloud platform. The default path is still the hardcoded URLs
183 of the shared S3 bucket of our GoToCloud platform.

184

185 C2-6) How long does it take to transfer data to the S3 bucket? Are you transferring from
186 Japan to an AWS region in the United States?

187

188 Response 2-6) We appreciate you bringing up another important aspect of AWS
189 regarding data transfer speed. We agree that this topic is of great interest to readers as
190 well. We have been constantly transferring cryo-EM data from KEK (Tsukuba, Japan) to
191 S3 buckets in the Tokyo AWS region. Also, we have less-regularly transferred cryo-EM
192 data from KEK to the US East (Northern Virginia) AWS Region for our developmental
193 purpose like the benchmark tests of the current work. To answer your questions, we
194 measured the times for the data transfers (three times each using 1.0 TB EMPIAR-
195 10581 dataset) as follow:

196 (1) The Lustre storage at KEK to a S3 bucket in the Tokyo AWS region:

197 Transfer time: 77.91 ± 1.39 minutes

198 Command line: `time aws --profile=kek903_kek-cs-scheme s3 sync`

199 `/lustre1/emdata/gotofly/gtc_benchmark/empiar-`

200 `10581/10581/data/EMD_0731_181122 s3://kek-cs-scheme-325321040781-`

201 `empiar10581-transfer-tokyo/EMD_0731_181122`

202 (2) The Lustre storage at KEK to a S3 bucket in the US East (Northern Virginia) AWS
203 region:

204 Transfer time: 235.18 ± 2.48 minutes

205 Command line: `time aws --profile=kek903_kek-cs-scheme s3 sync`

206 `/lustre1/emdata/gotofly/gtc_benchmark/empiar-`

207 `10581/10581/data/EMD_0731_181122 s3://kek-cs-scheme-325321040781-`

208 `empiar10581-transfer-us/EMD_0731_181122`

209 The data transfer times are always similar to the above and match with our experiences
210 over a few years of using AWS. To reflect these results, the following sentences were
211 added in line 178-187 of the Supplemental Information: “As examples, the averages of
212 times to transfer a 1.0 terabyte (TB) input dataset (EMPIAR-1058) for three times were
213 as follow: (1) 77.91 ± 1.39 minutes from KEK (Tsukuba, Japan) to S3 buckets in the

6

214 Tokyo AWS region using “`time aws s3 sync`” command on an on-premise computer at
215 KEK, (2) 235.18 ± 2.48 minutes from KEK to the US East (Northern Virginia) AWS
216 Region using the same method as (1), and (3) 90.26 ± 1.90 minutes from a S3 bucket to
217 the Amazon FSx Lustre filesystem within the US East (Northern Virginia) AWS region
218 using “`time cp -r`” command on the header node of the pcluster associated to the FSx. The
219 reported time (3) is averages of the execution time differences between the copy
220 command before preloading (148.43 ± 1.90 minutes) and after preloading (58.17 ± 0.00
221 minutes).” Here, the description of (3) is for Response 2-7.

222

223 C2-7) Also, how long does it take to transfer data from the S3 bucket to the Amazon FSx
224 Lustre filesystem?

225

226 Response 2-7) Thank you again for pointing out a vital data transfer for exploiting AWS
227 ParallelCluster. To answer your question, we measured the data transfer time from a S3
228 bucket to the Amazon FSx Lustre filesystem within the US East (Northern Virginia) AWS
229 region (three times using 1.0 TB EMPIAR-10581 dataset) and obtained the average data
230 transfer time of 90.26 ± 1.90 minutes. Again, the data transfer time is always similar to
231 this and match with our experiences.

232 The following command line was executed after deleting the source directory

233 `(/fsx/10581/Micrographs/)` from the AWS FSx Lustre, and then importing the metadata

234 of the files in this source directory from the S3 bucket associated to the pcluster. This

235 way, we ensured the `lfs hsm_state` of all files in the source directory are “released exists

236 archived”, meaning the preloading (i.e., data transfer time from a S3 bucket to the

237 Amazon FSx Lustre filesystem) is necessary to start using the files on the AWS FSx

238 Lustre.

239 `$ time cp -r /fsx/10581/Micrographs /fsx/20240810_measure_copy_before_preload`

240 The average data transfer time was 148.43 ± 1.90 minutes (executed three times). Then,
241 again executed the same copy command as follow.

242 `$ time cp -r /fsx/10581/Micrographs /fsx/20240810_measure_copy_after_preload`

243 The average data transfer time was 58.17 ± 0.00 minutes (executed three times). The

244 reported preloading times above is an average of the execution time differences between

245 copy before preloading and after preloading. Please refer Response 2-6 for the addition
246 of sentences associated this response.

247

248

249 Responses to Reviewer #3

7

250

251 C3-1) The manuscript concerns the computational requirements of Single Particle
252 Analysis for cryoEM. These can be high, and it is certainly true that structural biologists
253 often use inappropriate setups and hardware, and may do more processing than necessary.
254 Target resolution is also a key quantity, and achieving the highest resolution possible is
255 not always necessary for the biological question.

256

257 Response 3-1) Thank you for expressing your agreement regarding the target resolution.

258

259 C3-2) Many facilities and groups have sought to optimise the use of compute resources.
260 Some efforts have been published, most not. Relion has been benchmarked several times
261 before. Cianfrocco (ref 16) highlighted the benefits of spot instances on AWS in 2015.
262 The question is whether this article provides any additional insights or a practical way for
263 others to manage their resources.

264

265 Response 3-2) We thank the reviewer for pointing out the relevant research in this area.
266 The aim of our current work is on building a practical platform for computations required
267 in the cryo-EM field using the latest cloud computing technologies and currently available
268 services provided by AWS, rather than how to run RELION on AWS optimal way. As
269 reviewer #2 kindly summarized the essence of our paper concisely in "Overall
270 Impression", our approach can centralize the technical aspects of maintaining and
271 optimizing the data processing software, and with optimizing the use of AWS resources,
272 letting the users focus on processing their data. The architecture design and the utilization
273 of spot instances described in the Cianfrocco (ref 16) were nicely done using the
274 technologies and services available at that time. Here, we devised more recent
275 technologies and services provide by AWS. For example, we use AWS ParallelCluster
276 instead of STARcluster. This allows user to easily use the multiple latest AWS instance
277 types (e.g. G5, G4dn, C6i) within a single workflow of Cryo-EM SPA. Cianfrocco (ref
278 16) described the usage of a single EC2 instance type "r3.8xlarge" but there is no mention
279 about how to use the Amazon Machine Image (AMI) for the usage of multiple AWS EC2
280 instance types within a single workflow. Also, the FSx for Lustre is used in GoToCloud
281 as the storage whose access speed is significantly faster than elastic block storage (EBS)
282 used in Cianfrocco (ref 16). Therefore, to clarify our aim, we added the following
283 sentences in line 101-105 of the main text: "The aim is on establishing a practical
284 computational infrastructure for the cryo-EM field using the latest cloud computing
285 technologies through a publicly available service for both academia and industry. Our

8

286 approach can centralize the technical aspects of maintaining and optimizing the data
287 processing software, and with optimizing the use of AWS resources, letting the users
288 focus on processing their data."

289

290 C3-3) GoToCloud supports multi-node processing using a selection of common cryoEM
291 software via AWS pcluster. It manages the hardware selection, dependencies, and
292 cryoEM software installation. It also handles configuration with respect to AWS Regions.
293 It would appear to be a good solution for the authors, but it is unclear to me to what extent
294 the code or the article are useful to others.

295

296 Response 3-3) Thank you for pointing out an ambiguity in our writing regarding the
297 usefulness of GoToCloud. There are already users using the GoToCloud platform. We
298 have been receiving feedbacks from them saying that the GTC scripts and the online
299 documents including step-by-step instructions, which we have prepared for the
300 participants of the GoToCloud hands-on workshop, are easy to use and helpful. The users
301 became able to build GoToCloud platform instances for their own cryo-EM SPA projects
302 by themselves after a single online hands-on workshop with us. The online documents
303 were in Japanese but the translation to English is in progress (please also refer Response
304 1-1). To clarify the above point, we added the following sentences in line 168-176 of the
305 main text: "A manual procedure building a GoToCloud platform instance typically takes
306 approximately half a day to a full day even for the developers who have well familiarized
307 with the steps. With the GTC scripts and online documentations of step-by-step

308 instructions (see the Code availability section), this task can be completed in 30-60
309 minutes for users at any level of familiarity with the AWS system, making the GTC
310 scripts highly valuable. Moreover, since most of the time spent is just waiting for the
311 command executions, the actual time spent by the user is even shorter. In addition, by
312 carefully designing the specification of the GTC scripts, the possibility of user making
313 mistakes has been minimized.”

314

315 C3-4) Pre-installation of a range of executables is handled by a management account. It
316 was unclear at first whether the authors were providing a service to the community, but I
317 believe they are only providing the tools to do this. It is not even clear if it is being used
318 at the Photon Factory.

319

320 Response 3-4) Thank you for your interest in our service to the community in the
321 GoToCloud project. We have been providing not only the tool but also services to the

9

322 community by actively supporting the introduction of the GoToCloud platform to users'
323 own AWS accounts and the establishment of their supporting environments. Furthermore,
324 since July 2021, we have conducted several GoToCloud hands-on workshops for both
325 external users from academia and industry, with 1-4 participants each. As explained in
326 Response 3-3, there have already been both academic and industrial users utilizing the
327 GoToCloud platform for over a year, and internal use at our facility (i.e., the Photon
328 Factory) has also commenced.

329

330 C3-5) Maintaining up-to-date software is a headache for all facility and compute
331 managers. Software gets rapidly out-of-date, as shown by their usage of Relion 3 and 4
332 (version 5 is now out).

333

334 Response 3-5) We appreciate your understanding of the challenges faced by facility and
335 compute managers in keeping software up-to-date. The RELION5.0-beta (an online
336 release announcement on October 27, 2023) has been installed and can be used already
337 in the GoToCloud platform but this version is not set to default. Given that an unstable
338 version is likely to be superseded by a stable release soon, it is inefficient to invest
339 development efforts in it because repeating the optimization for each new software release
340 requires considerable effort. Therefore, our basic policy is to support only the latest
341 “stable” version as the default in the GoToCloud platform. Our plan is to make RELION-
342 5.0 default as soon as the stable release is announced, conduct large-scale benchmark
343 tests similar to ones provided in our current paper, and publish the results on our WEB
344 site. To clarify this, we added the following sentences in line 459-465 of the main text:
345 “RELION5.0-beta (an online release announcement in October 2023) has been also pre-
346 installed and can be used already in the GoToCloud platform but this version is not set to
347 default currently (RELION4.0 is current default). Given that an unstable version is likely
348 to be superseded by a stable release soon, it is inefficient to invest development efforts in
349 it because repeating the optimization for each new software release requires considerable
350 effort. Therefore, our basic policy is to support only the latest stable version as the default
351 in GoToCloud.”

352

353 C3-6) Their benchmarks show that the cost/performance can be optimised for each job
354 type. They choose a criterion of closest to the origin of a cost vs time graph, and I think
355 that is reasonable. They show that this is dependent on instance type, and will need to be
356 updated for as new hardware comes along. More urgently, I would assume that it is
357 dependent on the dataset. They have only benchmarked this calculation on a single

10

358 example (nitrite reductase) and so the variation in their conclusions is not tested. They do
359 not describe what would happen in normal usage. Does this benchmarking have to be
360 done on each new dataset in order to determine the best configuration? Is the user
361 expected to do this? Can anything be inferred from short runs (probably not as Relion
362 searches tend to be adaptive)?

363

364 Response 3-6) Thank you for pointing out that we missed to describe what would happen
365 in normal usage. In answer to the questions, similar benchmark tests do not need to be
366 performed for each dataset, and users are not expected to conduct benchmark tests to find
367 the optimal parallel settings for each of their datasets. Performing benchmarks on a test
368 case provides sufficient assurance that the overall performance for a workflow in general
369 use cases will achieve adequate level of optimality (please see below for the rationale).
370 We agree that this discussion is important, and so, the following paragraph is added in

371 line 383-400 of the main text. “An important consideration is the generality of proposed
372 optimal parallel computing settings for each RELION job type: here, we argue a rationale
373 for which benchmark tests do not need to be performed for each dataset and so a test case
374 using a single dataset should provide sufficient assurance that the overall performance of
375 a workflow in normal usage will achieve adequate level of optimality. Generally, for a
376 given job type (e.g. Refine3D), a “absolute” processing time and cost using a specific set
377 of parallel computing settings varies depending on dataset and processing step (where
378 Refine3D is used), because the computational loads are different while the computational
379 capability remains the same. On the other hand, “relative” relationship between two sets
380 of parallel computing settings should remain the same regardless of the computational
381 loads (e.g., “A” settings is always faster and cheaper than “B” settings), because, for each
382 relationship, the computational load is expected to remain relatively constant (possibly
383 with only minor fluctuations due to the use of random numbers by the algorithms) while
384 the computational capability of two sets of parallel computing settings are different. Thus,
385 by fixing the required computational load to a practical amount (not too small) and
386 optimizing the parallel computing settings (i.e., benchmarking for a specific dataset and
387 job), the optimal choice is expected to be nearly identical for a particular job type
388 regardless of dataset and processing stage at a practical level. “

389

390 C3-7) The test against resolution was not what I was expecting. They record the
391 computational requirements as the structure solution progresses through the different
392 stages, rather than testing against adjustable parameters (such as number of particles
393 picked, search parameters, etc). Possibly the only practical decision to be made is how

11

394 many cycles of particle polishing to do. Here, the conclusion is a bit strange. After
395 emphasizing the substantial cost of these additional cycles, they conclude it is worth doing
396 anyway. They suggest that the modest resolution gain will help model building, although
397 this is not actually tried. Overall, this test seems a bit pointless, unless they are going to
398 demonstrate some kind of stopping condition. Again, this test is done on a single
399 benchmark, so it is not clear how general it is.

400

401 Response 3-7) Thank you for pointing out that the motivation of the test against
402 resolution was not clearly explained. While optimization of adjustable parameters of each
403 algorithm type is important, our purpose of this test is different: it is a case study to give
404 an example of how the processing time and cost of a practical usage increases when one
405 tries to improve resolution by doing additional processes on the GoToCloud platform.
406 Our motivation was to find out if the relationships between the processing time and cost
407 relative to the resolution improvement are more like linear or exponential using a typical
408 SPA workflow. We believe that this example provides users a useful clue to decide a
409 minimum necessary resolution and for this up to which process stage they should employ
410 for a given research purpose. To clarify this point, we added the following sentences in
411 line 319-323 of the main text: “The motivation of this case study was to find how the
412 processing time and cost of a practical usage increases when one tries to improve
413 resolution by doing additional processes on the GoToCloud platform and if the
414 relationships between the processing time and cost relative to the resolution improvement
415 are more like linear or exponential using a typical SPA workflow.”

416 Thank you for your suggestion about Polish; we also think that a main practical
417 decision to be made is how many cycles of Polish to run. Therefore, we added the
418 following sentences in line 355-358 of the main text: “This result indicates that one
419 should consider omitting jobs in the later stage of the SPA processing if the resolution
420 sufficient for research purpose is achievable without them. Specifically, employing a
421 smaller number of Polish cycles is effective to reduce the processing time and cost.”
422 We appreciate your pointing out that the statement of “the modest resolution gain will
423 help model building” is misleading and not verified and we agree with this point.
424 Therefore, we removed the following sentences from the last paragraph of the
425 “Relationship between cost and resolution improvement in the Cryo-EM SPA
426 workflow” section in Result of the main text. “The appearances of the side chains
427 strongly affect the accuracy of atomic coordinate modeling and contribute to substantial
428 improvements in the modeling and associated statistical values. Therefore, in the

12

429 resolution range of around 2.0 Å, even if the improvement of the resolution is only in the
430 second decimal place, it should not be undervalued.”

431 Also, we agree that it is important to demonstrate some kind of stopping condition,

432 especially for the automation of the SPA workflow. However, this is not the scope of our
433 current work. Therefore, we modified the following sentence in line 413-416 of the main
434 text. "As the processing result of the streptavidin dataset indicated, an important issue for
435 the automation is to define the stopping condition of the workflow to achieve a high level
436 of practicality and should be addressed in this future study."

437 Because of a similar rationale explained in Response 3-6, performing a single
438 benchmark on a test case should provide a sufficient clue to readers about an expected
439 trend in the increases of total process time and cost relative to resolution improvement
440 for a typical processing workflow. However, again, validating this argument is beyond
441 the scope of this paper, it should be addressed in a future work for the full automation of
442 the cryo-EM SPA workflow.

443

444 C3-8) The core of GoToCloud seems to be a set of GTC configuration scripts. These are
445 available on github, and are a set of bash scripts. I could see no documentation there, and
446 I think it would be very difficult for anyone else to set up GoToCloud, let alone adapt it
447 or update it. There is some explanation in the Supplementary Methods which helps to
448 give the idea, but not enough for a 3rd party to adopt this solution.

449

450 Response 3-8) Thank you for reviewing the associated GitHub repository and for helping
451 us improve the accessibility and usability of our current work. As described in Response
452 1-1, we added two documentations to the GitHub: (1) README file with tutorial
453 information on how to use the GTC configuration scripts, and links to our web site
454 providing more elaborated version of the instructions, (2) License file to clearly indicate
455 under what license the software is released.

456

457 Some minor comments:

458 C3-9) 1. Figures 2 to 5 show the same plots for 4 different job types. I don't think this is
459 necessary and some space could be saved.

460

461 Response 3-9) Thank you for your feedback on figures 2 to 5. We can see why you would
462 make that suggestion since the visual impression of figures 2 to 5 share visual similarities.
463 However, we intended for all the main figures to clearly convey the key findings of our
464 benchmark tests. That is, the saturation points of the scalabilities (where the cost curve

13

465 becomes a vertical line), the minimum achievable cost (where the cost curve becomes a
466 horizontal line), and the optimal parallel computing settings are different among
467 Refine3D, Class3D, Class2D, and Polish. We would appreciate it if you could consider
468 keeping figures 2 to 5 the main figures, provided space permitting.

469

470 C3-10) 2. Figure 1a contains a typo - it should be UCSF Chimera(X).

471

472 Response 3-10) Thank you for finding a typo. It is fixed.

473

474 C3-11) In conclusion, although the work is technically correct, I don't think there is much
475 novelty in this article. The benchmarks and conclusions are specific to the two examples
476 studied, and there have been other benchmarks published. If it is meant to be a practical
477 solution, then there are many questions around how someone would deploy the
478 GoToCloud system, and how users would use it for their datasets. I cannot recommend
479 publication.

480

481 Response 3-11) We are grateful with your variable comments that allow us to remove
482 the ambiguities of our writing, especially ones related the novelty of our work and our
483 aim of each testing. As addressed in Response 3-2, the novelty of our current paper lies
484 in the newly proposed architectural design of the GoToCloud platform. This design
485 enables researchers to almost instantly access a powerful, ready-to-use computational
486 platform for cryo-EM SPA. Users can easily utilize multiple AWS instance types (e.g.,
487 G5, G4dn, C6i) within a single Cryo-EM SPA workflow. This feature facilitates high
488 efficiency for typical RELION workflows, which consist of multiple steps employing
489 different algorithms. Each of these algorithms requires specific configurations of CPUs,
490 GPUs, and memory for optimal performance. Additionally, the fast access speed of FSx
491 for Lustre enhances the performance of all algorithms. Importantly, benchmark tests have
492 led us to suggest optimal parallel computing settings, including the AWS EC2 instance
493 type, for each of RELION's Refine3D, Class3D, Class2D, and Polish steps. The
494 significant interest from over a dozen research groups who have requested the detailed
495 results of our benchmark tests underscores the need for a comprehensive resource within
496 the Cryo-EM SBDD community. To disseminate this valuable information and foster
497 further research, we would like to share our findings in this publication.

498 The deployment of the GoToCloud platform and how users can create and use the
499 platform instances for their datasets are addressed in Responses 1-1 and 3-8. Thus, it is
500 possible for readers to build and deploy their own GoToCloud platform. To help such

14

501 readers, we added a README file in our GoToCloud GitHub. The file contains the link
502 to the instruction of how to construct the shared EFS and S3, as well as how to install the
503 relevant software, such as RELION. The README file also describes the usage of the
504 GoToCloud scripts step-by-step to create and set up a GoToCloud platform instance on
505 user's AWS account.

506 We hope our responses clarify all of your questions and concerns.

15

Version 1:

Reviewer comments:

Reviewer #1

(Remarks to the Author)

Thank you for addressing my comments.

Reviewer #2

(Remarks to the Author)

Brief Summary

The authors provided very informative and complete responses to my questions.

Overall Impression

The improvements from optimizing Relion with CUDA_ARCH=86 versus CUDA_ARCH=75 are interesting. While the improvements were not uniform across the different processing steps, the speedups would save on the computing costs. Moreover, compiling software to run optimally on the hardware is the kind of optimization that would be helpful for general users who might not be familiar with GPU hardware.

Specific Comments

I only have specific comments about wording.

Line 223 - "the ground total cost" probably should be "the grand total cost"

Line 308 - "Core GPUs were tempted" probably should be "Core GPUs were attempted"

Supplemental, Line 300 - "more pronoun as the number of nodes" probably should be "more pronounced as the number of node"

Reviewer #3

(Remarks to the Author)

My main concern about the original manuscript was whether their solution was of general interest, or whether it was just another local solution. An additional concern was that only limited tests were presented for the benchmarks, so that it was unclear whether their proposed optimisations were generally applicable.

The authors have clarified that the GTC scripts have been used by others for over a year, as well as at the Photon Factory, which gives me some reassurance.

I highlighted the lack of documentation on the Github site, and Reviewer 1 also asked for a README file and a licence file. There is now a README file with a link to more detailed documentation (I haven't tried to follow the steps myself, but they look ok). There is a GPLv3 license file, although I can't find any source code headers specifying which files are covered by the license.

The authors have spelled out their policy for updating software versions. I understand the time required to maintain compatibility, but it is important that GoToCloud doesn't become obsolete in a year's time. I can see instructions for installing Relion5 under <https://sites.google.com/sbrj.jp/gotocloud-install/installation> but it is not clear whether any of the GTC scripts have been updated.

I am less convinced by the response on variation in datasets. The rebuttal states "users are not expected to conduct

benchmark tests to find the optimal parallel settings for each of their datasets". I believe the implication is that users should use the optimum setup from these benchmarks (Figs 2 – 5) for a given job type. So the question is, are these results transferable to other datasets? The authors argue that the relative performance of different compute configurations is constant, even if the absolute timings and costs change. This is just an assertion. As far as I can see, it has not been tested to see if it is true.

Anecdotally, Relion is not considered to scale well to multiple nodes, and probably most people use a single multi-GPU node. The vertical sections of the plots in Figs 2 – 5 are certainly believable. But the fine detail of whether to use 16@2x g4dn(8,96) or 8@1x g5(8,192) or 8@2x g5(4,48) for Refine3D seems a lot less certain. The Discussion again states "it is important to use the optimal settings for each combination of computer hardware and image processing job". If it is important, then I don't want to rely on benchmarks from a different system.

So a follow up question is what freedom does a user have? IIUC step 2 will create a pcluster with pre-defined EC2 instance types. Can users still select within these instance types and choose the number of nodes, threads, etc from within the Relion GUI accessed via NICE-DCV? In other words, are the benchmarks advice, or does GTC lock the user in?

In conclusion, I am mostly happy with the revised manuscript and it is suitable for publication. The extra detail is helpful, and the online documentation looks good. For the benchmarking, ideally I would have liked to see a second example from a different system (different resolution, different symmetry, multiple 3D classes, etc). If that is not possible, then the conclusions should be toned down. I am not convinced by the "rationale" in Response 3-6 and I don't think the additional text should be added. It would be better to simply say that the benchmarking is an example only.

With that one change, I can recommend publication.

Author Rebuttal letter:

1

2 Response to Reviewers

3

4 Responses to Reviewer #1 (Remarks to the Author):

5 C1-1) Thank you for addressing my comments.

6

7 Response 1-1) Thank you again for taking the time to review our manuscript. Your

8 feedback was very helpful in improving the accessibility and usability of our software.

9

10

11 Responses to Reviewer #2

12

13 Brief Summary

14 The authors provided very informative and complete responses to my questions.

15

16 Thank you for your positive feedback. We are pleased that our responses adequately

17 addressed your questions. Your expert comments on AWS were valuable in helping us

18 improve the quality of our paper.

19

20 Overall Impression

21 The improvements from optimizing Relion with CUDA_ARCH=86 versus

22 CUDA_ARCH=75 are interesting. While the improvements were not uniform across the

23 different processing steps, the speedups would save on the computing costs. Moreover,

24 compiling software to run optimally on the hardware is the kind of optimization that

25 would be helpful for general users who might not be familiar with GPU hardware.

26

27 We are grateful for your suggestion to optimize RELION with CUDA_ARCH=86 versus

28 CUDA_ARCH=75, which led to further improvements in processing speed and cost

29 performance.

30

31 Specific Comments:

32 C2-1) I only have specific comments about wording.

33 Line 223 - "the ground total cost" probably should be "the grand total cost"

34 Line 308 - "Core GPUs were tempted" probably should be "Core GPUs were attempted"

1

35 Supplemental, Line 300 - "more pronoun as the number of nodes" probably should be

36 "more pronounced as the number of node"

37

38 Response 2-1) Thank you for identifying the grammatical mistakes. All three mistakes
39 above have been corrected as you suggested.

40

41

42 Responses to Reviewer #3

43

44 C3-1) My main concern about the original manuscript was whether their solution was of
45 general interest, or whether it was just another local solution. An additional concern was
46 that only limited tests were presented for the benchmarks, so that it was unclear whether
47 their proposed optimisations were generally applicable.

48

49 The authors have clarified that the GTC scripts have been used by others for over a year,
50 as well as at the Photon Factory, which gives me some reassurance.

51

52 Response 3-1) We are glad to hear that the clarification regarding the use of the
53 GoToCloud platform and scripts by others provided you some reassurance. We hope that
54 our current paper conveys our message that our approach can be applied practically
55 beyond our immediate research environment, supporting its relevance and general interest.
56 Thank you for acknowledging this aspect of our work.

57

58 C3-2) I highlighted the lack of documentation on the Github site, and Reviewer 1 also
59 asked for a README file and a licence file. There is now a README file with a link to
60 more detailed documentation (I haven't tried to follow the steps myself, but they look ok).
61 There is a GPLv3 license file, although I can't find any source code headers specifying
62 which files are covered by the license.

63

64 Response 3-2) Thank you for taking the time to review our GitHub site as well. Following
65 your suggestion, we have added a header statement about the license and copyright to all
66 the source code files covered by the license.

67

68 C3-3) The authors have spelled out their policy for updating software versions. I
69 understand the time required to maintain compatibility, but it is important that
70 GoToCloud doesn't become obsolete in a year's time. I can see instructions for installing

2

71 Relion5 under <https://sites.google.com/sbrc.jp/gotocloud-install/installation> but it is not
72 clear whether any of the GTC scripts have been updated.

73

74 Response 3-3) We are grateful for your taking the time to review our online
75 documentation as well. To address your concern about an update of the GTC scripts in
76 conjunction with the installation of RELION5-beta, we designed the platform architecture
77 and GTC scripts to ensure that the scripts do not need to be updated whenever a new
78 version of RELION or other software is installed. As explained in the manuscript and on
79 the online instruction page, all supported versions of RELION and their associated
80 "modulefiles" are pre-installed and created on the shared EFS. The default version is
81 specified through the environment modules. We hope that the explanation clarifies your
82 concern.

83

84 C3-4) I am less convinced by the response on variation in datasets. The rebuttal states
85 "users are not expected to conduct benchmark tests to find the optimal parallel settings
86 for each of their datasets". I believe the implication is that users should use the optimum
87 setup from these benchmarks (Figs 2 – 5) for a given job type. So the question is, are
88 these results transferable to other datasets? The authors argue that the relative
89 performance of different compute configurations is constant, even if the absolute timings
90 and costs change. This is just an assertion. As far as I can see, it has not been tested to see
91 if it is true.

92

93 Response 3-4) Thank you for your valuable comments. While we agree that our rebuttal
94 statement is a logical deduction that needs to be tested, repeating benchmarks with
95 additional datasets was not feasible within the scope of this revision. Therefore, we
96 revised the additional text for Response 3-6 in our previous rebuttal to clarify that our
97 results are an illustrative example rather than a comprehensive validation. The revise also
98 acknowledges the importance of avoiding overly generalized conclusions based on a
99 single example. At the same time, we also address the C3-5 and C3-7 with this revised
100 paragraph (please refer Response 3-5 and Response 3-7). Accordingly, we modified the
101 corresponding paragraph in lines 382-398 of the main text, stressing in the last sentence
102 that the transferability of these results to other datasets needs to be tested in future work

103 (see below).
104
105 p. 17-18, lines 382-398.

3

106 “An important consideration is the general applicability of the proposed optimal parallel
107 computing settings for each RELION job type and/or other datasets. Our benchmarking
108 results, based on a single system and dataset, serve as an illustrative example specific to
109 the configuration used in this study. While these results offer valuable insights into
110 optimizing parallel computing settings on the GoToCloud platform, they are primarily
111 intended to help users customize these settings according to their research priorities, as
112 discussed above. Unlike on-premise computing resources, which are often not publicly
113 accessible and likely offer limited availability to external users, the AWS cloud
114 computing service is publicly available. This accessibility makes it easier for researchers
115 in both academia and industry to optimize parallel computing settings on the same
116 platform according to their needs by referring our benchmark results. Nevertheless, our
117 benchmarking results should be viewed as an example rather than a comprehensive
118 validation. Therefore, outside the GoToCloud platform, we encourage readers to conduct
119 their own benchmarks on their specific systems, to achieve optimal performance.
120 Additionally, the transferability of these results to other datasets is another important
121 consideration, and this should be explored in future work to fully automate the cryo-EM
122 SPA workflow.”

123

124 C3-5) Anecdotal, Relion is not considered to scale well to multiple nodes, and probably
125 most people use a single multi-GPU node. The vertical sections of the plots in Figs 2 – 5
126 are certainly believable. But the fine detail of whether to use 16@2x g4dn(8,96) or 8@1x
127 g5(8,192) or 8@2x g5(4,48) for Refine3D seems a lot less certain. The Discussion again
128 states “it is important to use the optimal settings for each combination of computer
129 hardware and image processing job”. If it is important, then I don’t want to rely on
130 benchmarks from a different system.

131

132 Response 3-5) Thank you for pointing out the ambiguity regarding the important aspect
133 of the applicable range of our benchmark tests. We agree that users should not rely solely
134 on benchmarks from a different system but should consider them as a guideline or starting
135 point for finding optimal parallel computing settings on their own hardware. Therefore,
136 we clarified that our results are an illustrative example rather than a comprehensive
137 validation and overly generalized conclusions based on this single example should be
138 avoided. This point is also addressed in Response 3-4.

139

140 C3-6) So a follow up question is what freedom does a user have? IIUC step 2 will create
141 a pcluster with pre-defined EC2 instance types. Can users still select within these instance

4

142 types and choose the number of nodes, threads, etc from within the Relion GUI accessed
143 via NICE-DCV? In other words, are the benchmarks advice, or does GTC lock the user
144 in?

145

146 Response 3-6) Thank you for your interest in the GoToCloud platform. Through the
147 RELION GUI application in the NICE-DCV WEB page of GoToCloud, users can choose
148 a Slurm partition, which is associated with a specific EC instance type, for each RELION
149 job from a set of available partitions and specify the number of nodes and threads. To
150 depict this, we added Supplemental Fig. S2b that shows a terminal window with the list
151 of available partitions printed by the Slurm “sinfo” command. Finally, this is why we
152 provide the details of the benchmark results as guidance or a starting point of the EC
153 instance type selection for users.

154

155 C3-7) In conclusion, I am mostly happy with the revised manuscript and it is suitable for
156 publication. The extra detail is helpful, and the online documentation looks good. For the
157 benchmarking, ideally I would have liked to see a second example from a different system
158 (different resolution, different symmetry, multiple 3D classes, etc). If that is not possible,
159 then the conclusions should be toned down. I am not convinced by the “rationale” in
160 Response 3-6 and I don’t think the additional text should be added. It would be better to
161 simply say that the benchmarking is an example only.

162

163 With that one change, I can recommend publication.

164

165 Response 3-7) Thank you for your thorough review and for your positive feedback on
166 the revised manuscript and the online documentation. Regarding your feedback on
167 Response 3-6 of our previous rebuttal, we appreciate your input and agree that the
168 additional text should be removed. We modified the text added in Response 3-6 of our
169 previous rebuttal based on your suggestion. This point is also addressed in Response 3-
170 4.

171

172 We hope these changes address your concerns, and we appreciate your constructive
173 suggestions.

174

5

Open Access This Peer Review File is licensed under a Creative Commons Attribution 4.0 International License, which permits use, sharing, adaptation, distribution and reproduction in any medium or format, as long as you give appropriate credit to the original author(s) and the source, provide a link to the Creative Commons license, and indicate if changes were made.

In cases where reviewers are anonymous, credit should be given to 'Anonymous Referee' and the source.

The images or other third party material in this Peer Review File are included in the article's Creative Commons license, unless indicated otherwise in a credit line to the material. If material is not included in the article's Creative Commons license and your intended use is not permitted by statutory regulation or exceeds the permitted use, you will need to obtain permission directly from the copyright holder.

To view a copy of this license, visit <https://creativecommons.org/licenses/by/4.0/>