



Singapore Informatics League 2026

Post-Contest Report

September 2026

Supported by:



Sponsored by:



Contents

SIL 2026 at a Glance	ii
1 Introduction	1
2 Participant Demographics	1
2.1 Comparison with SIL 2025	2
3 Online Contest	3
3.1 Batch Task Analysis	3
3.2 Optimisation Task Analysis	5
3.2.1 Scoring Formula	5
3.2.2 Engagement and Score Progression	5
3.3 Issues Reported During the Contest	5
3.4 Post-Contest Verification Review	7
3.5 Award Boundary Calculation	8
4 Final Contest	9
4.1 Event Overview	9
4.2 Batch Task Analysis	9
4.3 Optimisation Task Analysis	10
4.3.1 Scoring Formula	10
4.3.2 Engagement and Score Progression	11
4.4 Issues Reported During the Contest	11
4.5 Award Boundary Calculation	12
4.6 Closing Ceremony	13
5 Post-Contest Survey Results	15
5.1 Online Contest	15
5.2 Final Contest	15
5.2.1 Overall Reception	15
5.2.2 Issues Highlighted	16
6 Evaluating SIL's Objectives	17
6.1 Lowering Barriers to Entry	17
6.2 SIL as a Stepping Stone towards NOI	19
6.3 SIL as a Team-based Competition	20
6.4 Representation and Access to the Final Contest	21
7 Conclusion	23
A Participant Demographics	24
B Batch Task Statistics	25
C Optimisation Statistics	33
D Survey and Qualitative Feedback	34
E Task Topic Classification	37

SIL 2026 at a Glance

Participant Demographics

757 participants forming 177 teams from 70 schools registered for SIL 2026, an increase of 137 participants and 27 teams from 2025. 64.1% had not previously participated in the NOI, 51.7% had no formal computing qualification, and 53.8% reported a self-assigned skill level of 1 (Beginner).

Online Contest (5 July 2026)

The contest comprised 25 batch tasks and one optimisation task. 142 teams appeared in the official standings, excluding 27 disqualifications and 8 inactive teams. 43 teams received a Silver Award and 43 a Bronze Award. 34 teams received invitations to the Final Contest; 32 accepted and competed, including 8 wild-card teams.

Contest Integrity

46 teams were reviewed for potential rule violations; 27 were disqualified, and 6 individuals were banned from future editions. All participants were required to record their screens, and these recordings were essential to our review process.

Final Contest (15 August 2026, NUS School of Computing)

Our inaugural Final Contest was a full-day event including a four-hour contest, a task analysis session, a talk and question-and-answer session on studying Computing at NUS, and a closing ceremony. 143 of the 148 registered finalists attended, representing 27 institutions. The contest consisted of 10 batch tasks and one optimisation task, and 16 teams received a Gold Award.

Post-Contest Survey Results

Survey results were generally positive: the contest platform, the tasks, and the organisation of the Final Contest were all rated highly. The main reservations concerned the overall accessibility of the task sets and engagement with the Online Contest optimisation task, both of which we address in our priorities for 2027.

Priorities for SIL 2027

Our priorities for SIL 2027 fall into four areas: broadening access to the Final Contest, maintaining the integrity measures that proved effective this year, making the harder parts of the task set more approachable, and refining the on-site event experience. The specific commitments behind each of these appear in the relevant sections of this report.

1 Introduction

The **Singapore Informatics League (SIL)** is a team-based informatics (computer science) competition open to students studying in secondary schools or pre-tertiary institutions in Singapore. It emphasises teamwork and collaboration, and is intended to serve as a beginner-friendly introduction to competitive programming contests. A key objective of SIL is to promote computer science education among students of diverse backgrounds, such as those from traditionally underrepresented schools and genders.

The second edition of SIL comprised two contests. The Online Contest was held on 5 July 2026, and the inaugural Final Contest was held on 15 August 2026 at the NUS School of Computing.

SIL 2026 was, again, supported by the **Centre for Nurturing Computing Excellence (CeNCE)**, which shares our commitment to developing computing excellence among both pre-university and university-level students. This year, we were fortunate to have the financial support of **Jane Street**, a quantitative trading firm and liquidity provider with a unique focus on technology and collaborative problem solving. This edition of SIL would not have been possible without the support of CeNCE and Jane Street.

This report presents key contest statistics and outlines the Organising Team’s considerations in addressing issues that arose during and after the contests.

2 Participant Demographics

757 participants forming 177 teams registered for SIL 2026. 70 different schools were represented. These included secondary schools, junior colleges, polytechnics, and international schools.

School distribution of all participants

All 757 registered participants from 70 distinct institutions; a school is counted under each type it appears under

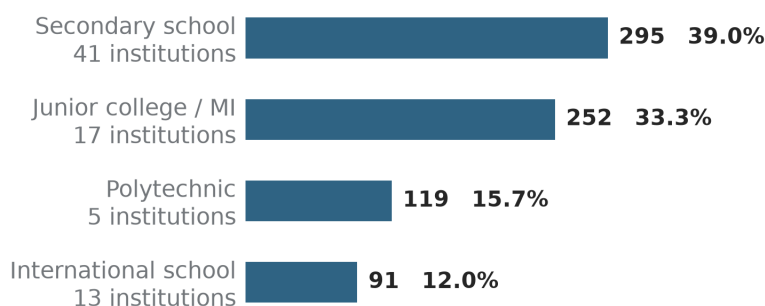


Figure 1: Institutional reach among the participants.

More than half did not have any formal computing qualifications, and more than half reported a self-assigned skill level of 1 (Beginner). Nearly two-thirds did not previously participate in the National Olympiad in Informatics (NOI). The number of newcomers increased from the 2025 edition, indicating that interest in SIL among beginners is growing.

Newcomers among participants

All 757 registered participants; categories are independent and overlap

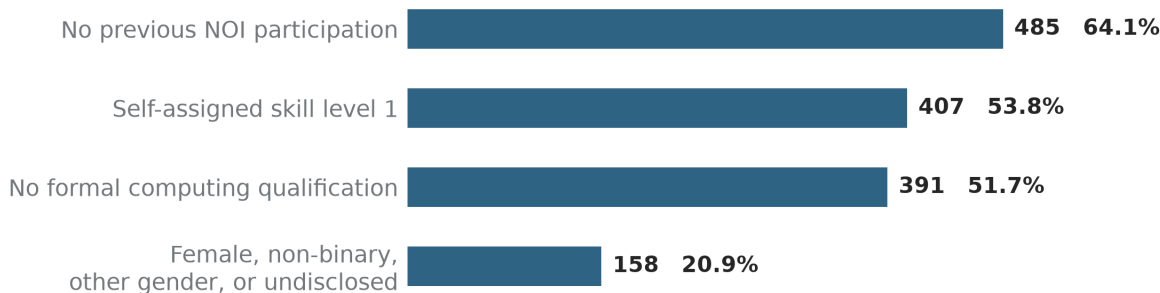


Figure 2: Accessibility indicators among the participants.

409 participants came from schools with no NOI Gold medallist in the five most recent editions. They formed 54.0% of registrations but only 20.3% of finalists; this is discussed further in Section 6.4.

2.1 Comparison with SIL 2025

SIL 2026 saw an increase of 27 registered teams and 137 participants. Figure 3 provides a visualisation of participant retention across the two editions. Of the 25 SIL 2025 participants who were no longer eligible (due to having graduated from a pre-tertiary institution), 7 joined the [SIL 2026 Organising Team](#).

Figure 4 compares selected statistics across editions. Notably, female, non-binary, other-gender, and undisclosed-gender participants formed 20.9% of registrations, down 8.8 percentage points from 2025.

The full 2026 demographic breakdown and the 2025 and 2026 comparison can be found in Appendix A.

Participant overlap between SIL 2025 and 2026

25 of the 414 people who didn't return to SIL 2026 were ineligible, and 7 joined the Organising Team

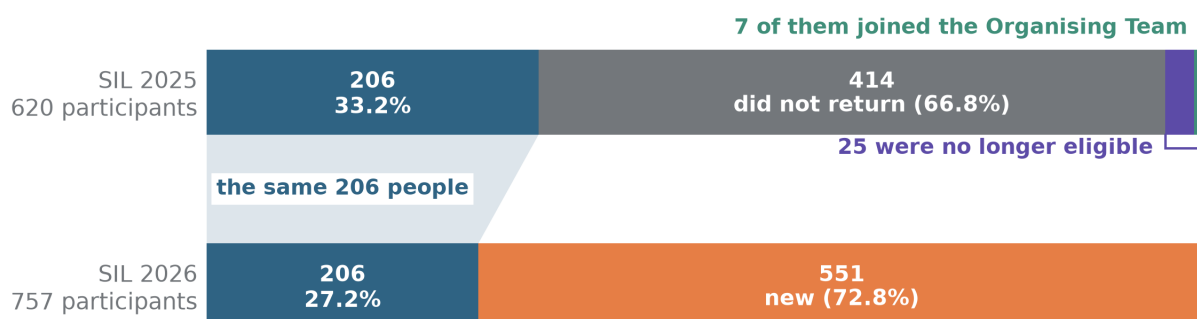
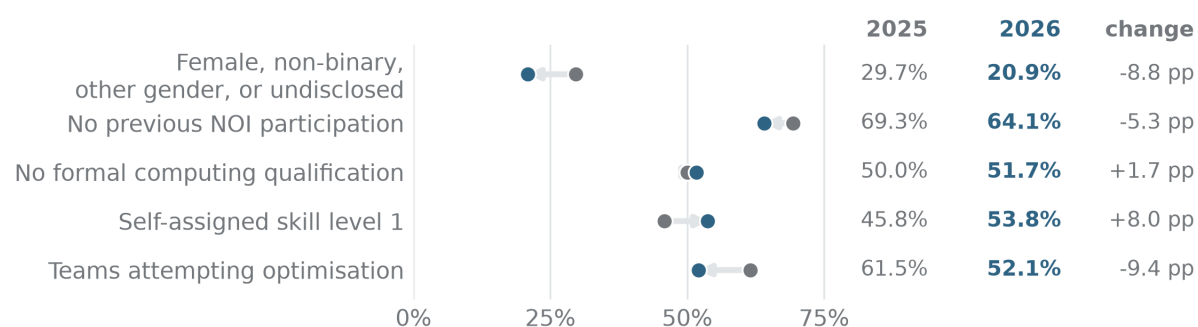


Figure 3: Participant retention and growth from SIL 2025 to SIL 2026.

Accessibility indicators moved in different directions

Arrows run from the 2025 value to the 2026 value. Both demographic cohorts are registration cohorts; optimisation counts teams in the official results in both years.



The two editions differ in format, so the pairs remain indicative rather than controlled. The demographic comparison table gives the same registration measures as counts.

Figure 4: Selected accessibility measures in 2025 and 2026. Population and format differences should be considered when interpreting the changes.

3 Online Contest

3.1 Batch Task Analysis

The Online Contest featured a total of 25 batch tasks. Each task was assigned a difficulty “level” in accordance with the [SIL Syllabus](#), which defines four levels. Level 1 tasks require basic mathematical thinking and can be solved without programming; Level 2 tasks require introductory programming knowledge; Level 3 tasks require elementary knowledge of algorithms and data structures; and Level 4 tasks are set at a difficulty comparable to the NOI. The Online Contest used Levels 1 to 3, with 8, 8, and 9 tasks in each level respectively. All tasks were original;

we created them specifically for SIL 2026. The **task statements** and their **solutions** have been published on our website.

The Scientific Committee reviewed 80 task proposals over several months. Some of these proposals were authored by the Scientific Committee itself, while others were submitted through our informal “Call For Tasks”. We evaluated each proposal for pedagogical value and quality, and selected the task set while ensuring a balanced distribution of topics and difficulty levels. The primary and secondary topics of each Level 2 and Level 3 task are listed in Appendix E.

Batch task solve rates

Cohorts use prior NOI experience: none (All-new), below Silver (Some prior), or at least one Silver or Gold medallist (Advanced).

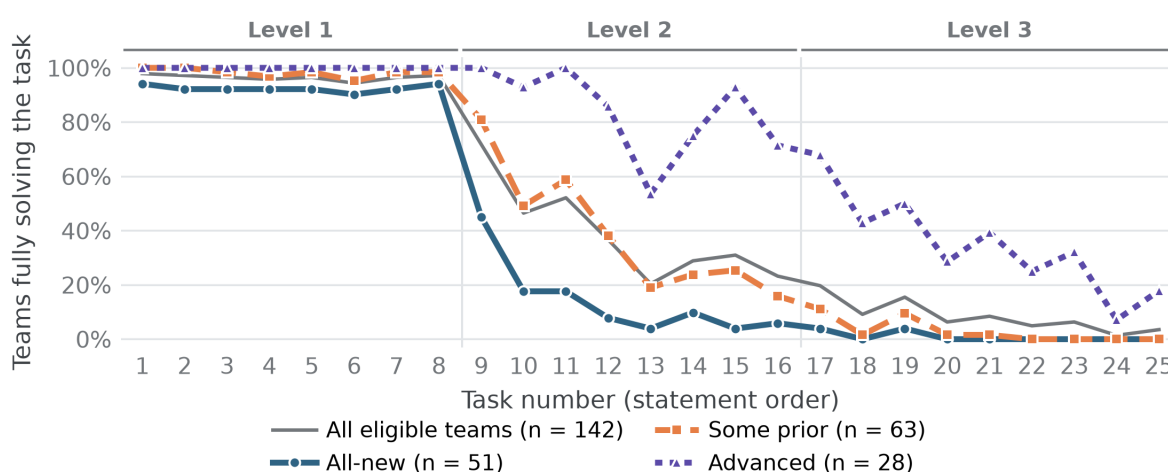


Figure 5: Full-solve rate for each Online Contest batch task, with teams grouped by prior NOI experience.

We observe that the first eight tasks were approachable for nearly all teams, with full-solve rates all exceeding 90%. This indicates that most teams were able to stay engaged with the Level 1 tasks.

We then see a steep decrease in full-solve rate among Level 2 tasks. This year’s set of Level 2 tasks was harder than last year’s: while all Level 2 tasks in 2025 had a greater than 40% full-solve rate, 5 out of 8 of 2026’s Level 2 tasks had a rate of less than 40%. Level 3 tasks were effective at distinguishing the highest-performing teams, with six of their nine tasks fully solved by at most 9% of teams.

Many teams attempted to earn points on easier subtasks. Among Level 2 and Level 3 tasks, the median conversion from Subtask 1 to a full solve was 62% and 64% respectively. This suggests that the subtasks were approachable, allowing teams to score points on them even if they could not fully solve the task. Appendix B gives the complete task and subtask statistics.

3.2 Optimisation Task Analysis

The optimisation task, Watercup Robot, required participants to strategically guide a robot to collect and deliver as much water as possible while managing limited time and water lost during movement. Of the 142 eligible teams, 74 attempted at least one optimisation subtask (52.1%), down from 61.5% in 2025. We will introduce the optimisation task more clearly in future editions.

3.2.1 Scoring Formula

Each subtask was worth 15 points. For a team of rank r among $N = 177$ registered teams (the published rule counts all registered teams, including those subsequently disqualified), with raw score x and best raw score b , the published rule was

$$10 \left(1 - \frac{r-1}{N} \right)^{1.25} + 5 \left(1 - \left(1 - \frac{x}{b} \right)^{0.4} \right).$$

The first term rewards rank, and the second rewards closeness to the best construction. Each subtask is scored separately, and (as a special case) a team without a valid submission received zero. Under this rule, any valid submission earned points.

3.2.2 Engagement and Score Progression

Most teams that attempted Subtask 1 attempted Subtask 3 as well. Attempts fell from 74 teams on Subtask 1 to 55 on Subtask 3. Subtask 1 was solved by half of the teams attempting it, with 37 of the 74 teams reaching the best raw score.

For Subtask 2, scores plateaued relatively early, with contestants reaching a score of 54,062 around the one-hour mark. The score is likely at or close to the optimum; the Scientific Committee’s reference solution reached the same value. In Subtask 3, one team stood out with a score of 94,254, exceeding the Scientific Committee’s 73,571.

Subtask 3 saw the largest gains in score over the duration of the contest: its record rose from 22,142 to 94,254, an improvement of 325.7%. Major improvements occurred in the first hour, and smaller gains continued around the scoreboard freeze, with the final record set after it.

3.3 Issues Reported During the Contest

We were informed of a minor typo in the example explanation for Task 21: Spinning Table. We fixed the issue promptly and alerted participants to the matter.

At approximately one hour into the contest, we discovered that the test case for Task 13: Tokens was incorrect. The correct test case was re-uploaded at 1008H. We informed teams about

Online Contest normalised optimisation scores, by subtask

Teams that attempted each subtask, on the published 15-point scale

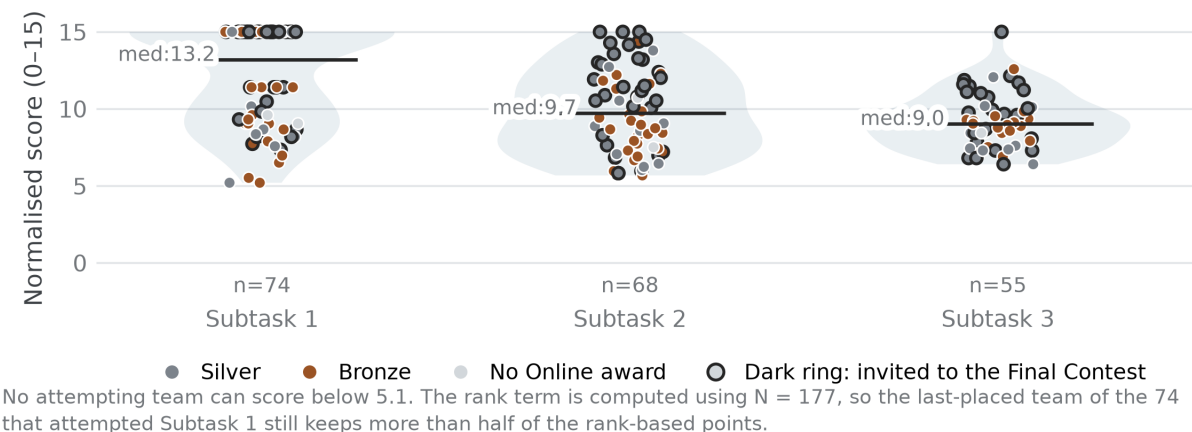
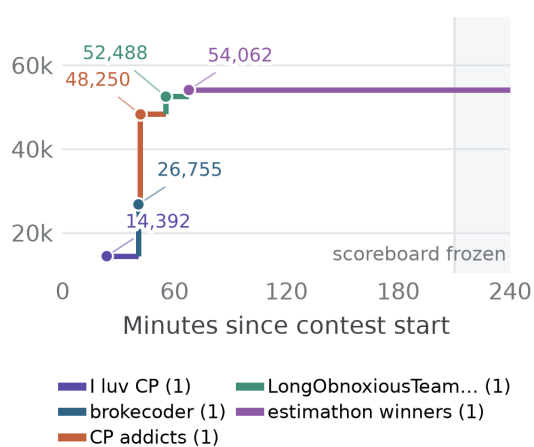


Figure 6: Normalised optimisation scores for teams that attempted each subtask. Raw-score distributions are in Appendix C.

Online Contest: Subtask 2

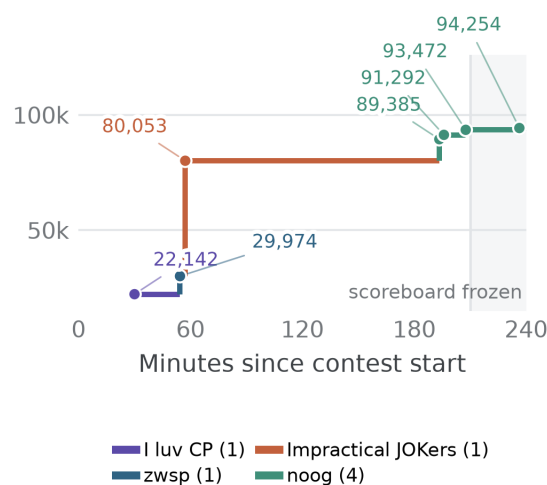
Leading raw score, 14,392 to 54,062, a 276% gain



(a) Subtask 2.

Online Contest: Subtask 3

Leading raw score, 22,142 to 94,254, a 326% gain



(b) Subtask 3.

Figure 7: Every improvement to the leading raw score on the two Online Contest optimisation subtasks.

the change through the contest platform, but some participants did not notice the announcement and continued to work with the incorrect version of the test case. We will make future announcements more visible to prevent a repeat of this.

At 1015H, we announced that Subtask 3 of Task 19: Dictator was worth 5 points rather than the 8 points stated.

During the contest, several participants reported difficulties related to handling large input files, including being unable to correctly read the provided input into their programs. Additionally, some participants raised concerns about slow program execution or insufficient memory on their workstations when running their solutions.

These issues fall within the scope of participant responsibility, and the Organising Team was therefore unable to provide technical assistance during the contest or accept appeals on these grounds. All contest tasks are designed with the guarantee that a correct solution exists that runs efficiently within reasonable time and memory limits (15 seconds, 1024 MB) on standard hardware. Consequently, appeals based on perceived unfairness arising from differences in participant workstation performance cannot be accepted.

This is an issue that we identified in SIL 2025. This year, we published a [Beginner's Handbook](#) that explains how to circumvent these technical difficulties. Participants are highly encouraged to read the Handbook for further guidance.

3.4 Post-Contest Verification Review

After conducting a review of participants' submissions and contest activity, we identified 46 teams that potentially violated the rules by using Generative Artificial Intelligence tools or seeking outside assistance beyond their registered team.

Out of these 46 teams:

- 19 were cleared of suspicion of wrongdoing.
- 12 teams did not respond to our request for clarification, which we considered an implicit admission of wrongdoing.
- 3 teams admitted to cheating when prompted.
- 5 had clear and undeniable signs of rule violation in their screen recordings or submitted code files.
- 3 produced inconsistent evidence in their appeals, leading to disqualification.
- 4 were invited to an online meeting to resolve our concerns about their contest activity. All 4 were unable to produce sufficient evidence to justify their contest behaviour, and were disqualified.

As a result, 27 teams were disqualified from the contest. This represents 15.3% of the 177 registered teams. Together with 8 teams that registered but did not compete, this leaves 142 teams in the official Online Contest results; unless otherwise stated, all Online Contest statistics in this report are based on these 142 teams. Additionally, we deemed 6 individuals to have been

deceptive about their contest behaviour, and have banned them from participating in future editions of SIL.

This investigation process took two to three weeks, required much effort from the Organising Team, and delayed the release of the results, but it was essential to maintain the integrity of the contest. Upholding these standards of fairness is crucial not only for SIL 2026, but for future editions as well. We will streamline and strengthen this review process in future editions.

This year, all participants were required to record their screens for the entire duration of the contest. These recordings, together with the submitted code files, provided most of the evidence in the review, and we found the requirement to be very effective. We will retain it in future editions.

We would like to remind all participants that SIL is designed to be a relaxed, beginner-friendly introduction to computer science competitions. The goal of the contest is for participants to enjoy learning, challenge themselves, and collaborate together, not to win at all costs.

3.5 Award Boundary Calculation

In accordance with the [rules](#), the top 60% of teams won at least a Bronze Award and the top 30% of teams won a Silver Award.

The published Silver band ended at rank 43 (123.85 points), while Bronze ended at rank 86 (80.90 points). The initial direct invitations to the Final Contest went to the top 24 teams; the rank-24 threshold was 147.49 points. 8 other teams qualified through wild-card entry.

Online Contest total-score distribution

The teams that received a Finals invitation are ringed; the dashed rules mark the award boundaries and the Top 24 cutoff.

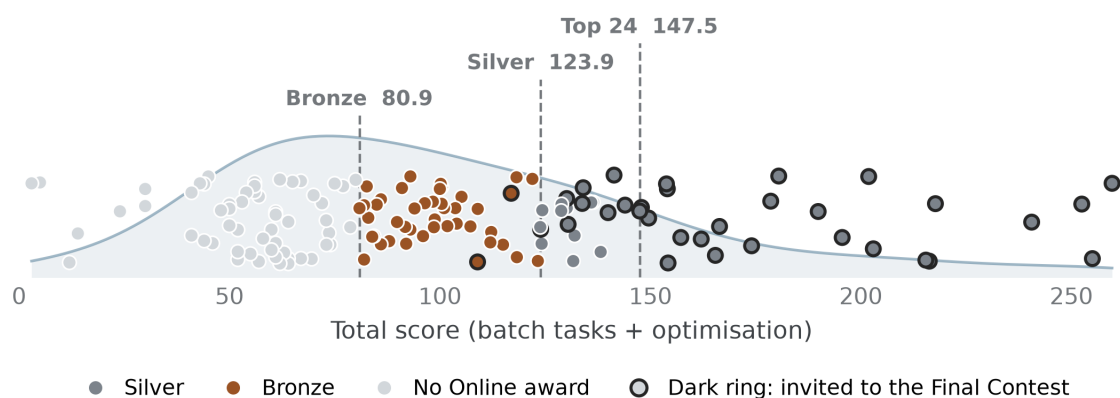


Figure 8: Online total-score distribution for 142 eligible teams. Fill colour shows Silver Award, Bronze Award and no-award results; a dark ring marks all 34 Final Contest invitation recipients, including the two teams that declined.

4 Final Contest

4.1 Event Overview

The inaugural Final Contest was held at the NUS School of Computing COM3 Multi-purpose Hall (MPH). 32 teams participated, representing 27 pre-tertiary institutions (after secondary and junior college sections of the same institution were merged). Of the 148 registered finalists, 143 attended the event.

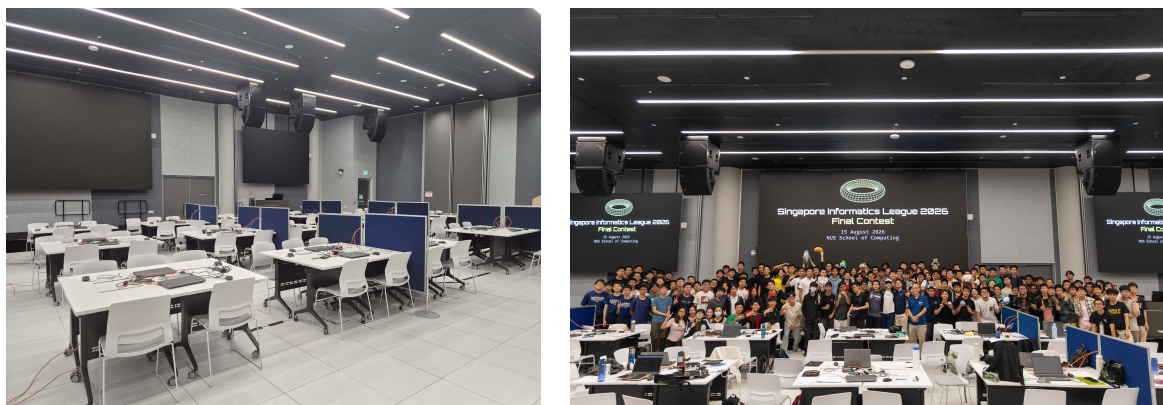


Figure 9: The COM3 MPH contest venue and the closing group photograph.

The event included a four-hour contest, task analysis by the Scientific Committee, a talk and question-and-answer session about studying Computing at NUS, and the closing ceremony. The admissions talk was given by Associate Professor Wong Weng Fai, Associate Dean of Undergraduate Admissions and Deputy Head of the Computer Science Department.

4.2 Batch Task Analysis

The Final Contest consisted of 10 batch tasks, each worth 15 points across three subtasks. We hid the task difficulty levels and presented the tasks in alphabetical order, rather than from easiest to hardest.

- Level 2: 1. 6 ate 8, 4. Reveille, 6. Snowwabbits
- Level 3: 2. Dog Walk, 3. Gremlin Nob, 9. Topsy Turvy
- Level 4: 5. Slimes, 7. Spring Sunlight, 8. StringStringString, 10. Xiaoyang Ranking

The three Level 2 tasks were fully solved by between 78% and 81% of teams. The Level 3 tasks were more effective at distinguishing the stronger teams, with full-solve rates falling to 38%, 16%, and 6%. At Level 4, three tasks were solved by only 3–6% of teams, while Slimes

was solved by none. It received 30 rejected submissions, suggesting that it was too difficult to provide useful separation among teams.

Full-solve rates in the Final Contest

32 finalist teams, grouped by prior NOI experience. The All-new cohort has one team, so its solve rates are either 0% or 100%.

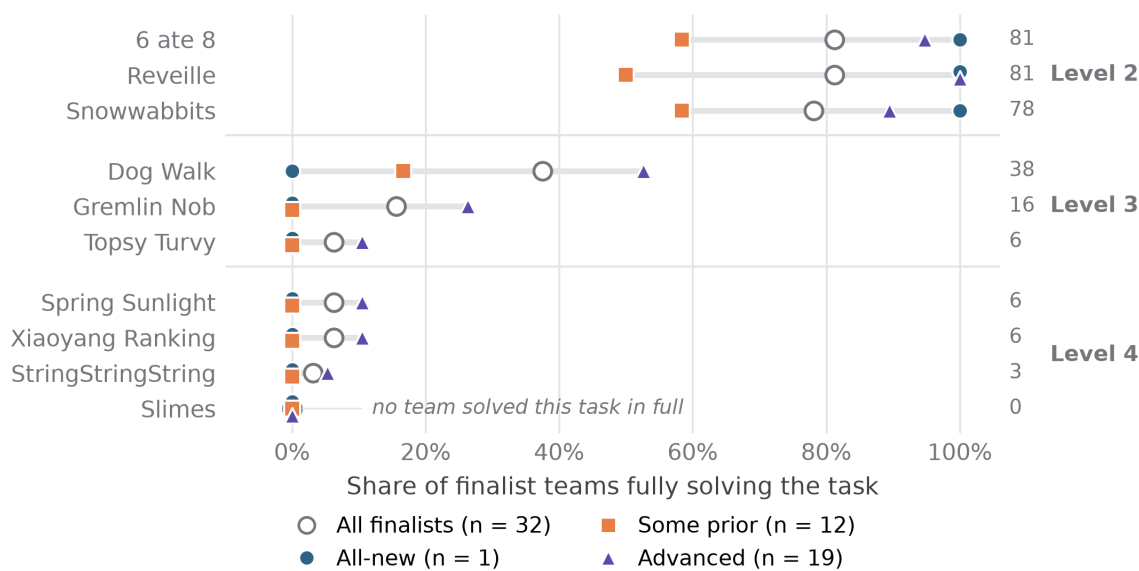


Figure 10: Final Contest full-solve rates for all 32 finalist teams, and for teams grouped by prior NOI experience.

4.3 Optimisation Task Analysis

The optimisation task, Civilisation, tasked participants with strategically building roads to connect mines, gather resources, and maximise gold production within a limited number of turns.

4.3.1 Scoring Formula

The number of points allocated to each subtask and the normalisation formula were identical to the Online Contest, with the only difference being $N = 32$ instead of $N = 177$.

4.3.2 Engagement and Score Progression

All 32 teams attempted the optimisation task. Refer to Appendix C for detailed statistics. Based on informal feedback from participants, the Final Contest optimisation task was easier to approach using the provided visualiser, which lowered the skill floor required to obtain a good score.

Final Contest normalised optimisation scores, by subtask

Teams that attempted each subtask, on the published 15-point scale

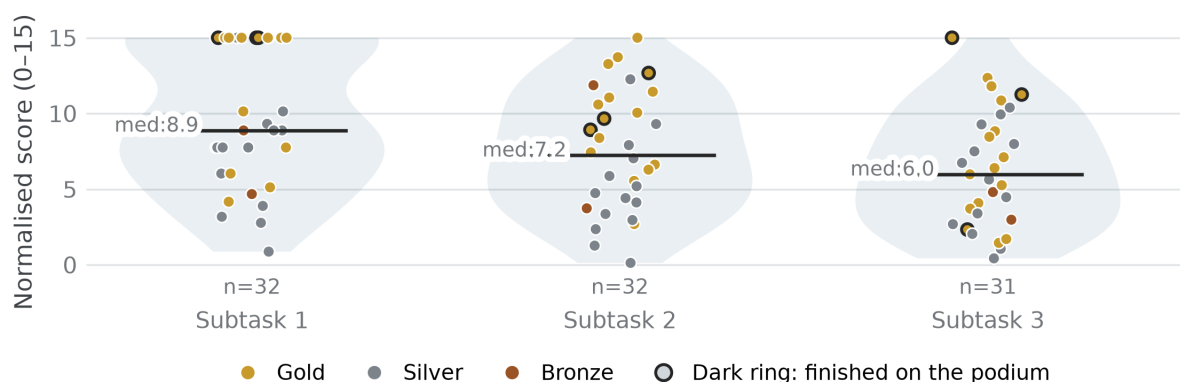


Figure 11: Normalised Final Contest optimisation scores for teams that attempted each subtask.

The record for Subtask 2 rose from 11,364 at the 24-minute mark to 12,420 at the 205-minute mark, an improvement of 9.3%. The record for Subtask 3 reached 1,360,701 at the 195-minute mark. Both final records were set after the scoreboard freeze and were revealed during the Resolver. The median raw scores were 10,937.5 for Subtask 2 and 781,143 for Subtask 3, against best scores of 12,420 and 1,360,701 respectively, indicating that Subtask 3 provided considerably greater separation between teams.

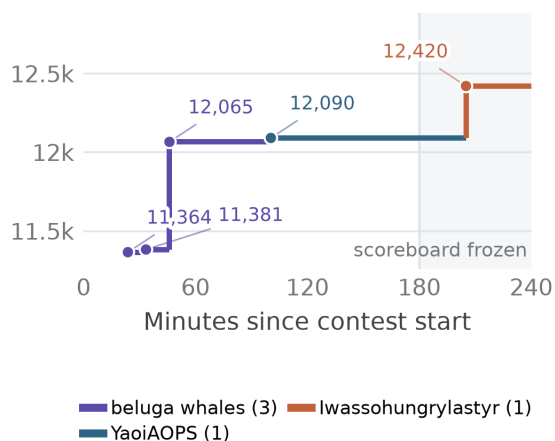
4.4 Issues Reported During the Contest

Near the beginning of the contest, a technical issue resulted in some teams being unable to access the visualiser for the optimisation task as a pop-up on the contest platform. As the Technical Committee worked to fix the issue, we informed teams that they would be able to access the visualiser if they opened it in a new tab. The issue was resolved within 20 minutes.

Approximately 30 minutes into the contest, it was discovered that the answers for Subtasks 2 and 3 of Task 1 (6 ate 8) were incorrect. Several teams had submitted correct answers that were incorrectly deemed as wrong answers. The issue was resolved within 10 minutes, and verbal and written announcements were made to inform teams of the mistake. We allowed teams to submit appeals if they felt that they were affected by the issue; of the 11 teams that appealed, 10 had their scores adjusted accordingly.

Final Contest: Subtask 2

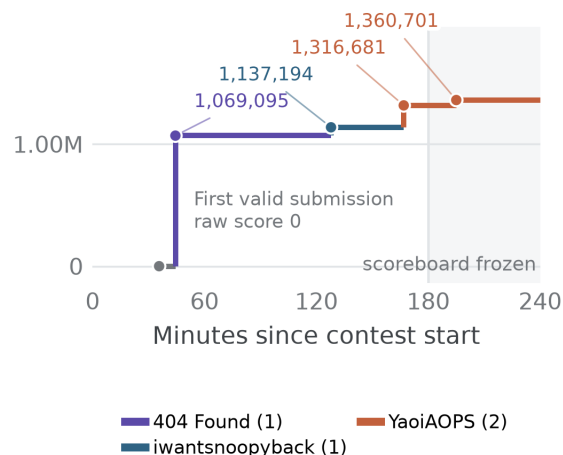
Leading raw score, 11,364 to 12,420, a 9% gain



(a) Subtask 2.

Final Contest: Subtask 3

Leading raw score, 0 to 1,360,701



(b) Subtask 3.

Figure 12: Every improvement to the leading raw score on the two Final Contest optimisation subtasks.

Approximately one hour into the contest, a team informed us that the constraints for Subtask 2 of Task 4 (Reveille) were incorrectly stated. We updated the statement promptly, and announced verbally and on the platform to inform teams of the correction.

Errors in task statements, scoring information and test data occurred across the two contests; we will tighten our technical checks to prevent similar errors.

4.5 Award Boundary Calculation

The top 16 teams earned Gold based solely on their Final Contest performance, while the remaining 16 kept the award they had earned in the Online Contest. The top three teams also received Amazon vouchers as prizes: \$500 for first, \$300 for second, and \$200 for third.

Final Contest total-score distribution

The teams that finished on the podium are ringed; the Gold and Top 3 boundaries are marked.

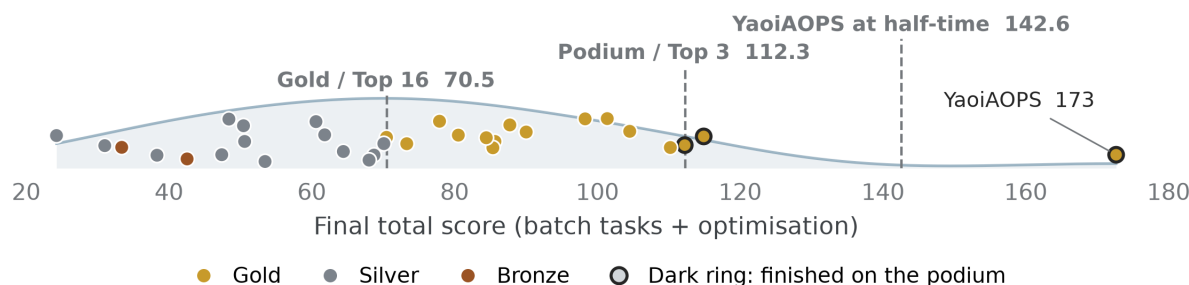


Figure 13: Final Contest total-score distribution for the 32 finalist teams. The Gold Award and podium boundaries are marked.

4.6 Closing Ceremony

The scoreboard was frozen in the final hour of the contest, with the submissions made during that time period, along with the optimisation scores, revealed during the closing ceremony in an International Collegiate Programming Contest (ICPC)-style Resolver segment.

The Resolver began from the lower end of the frozen standings, progressively resolving each team's pending submissions, with successful submissions causing teams to move upwards before the process continued. There were 349 submissions after the freeze, and 22 teams increased their batch scores. Between the frozen batch-only scoreboard and the official final standings, which also included optimisation scores, 29 of the 32 teams changed rank.

Selected rank changes after the scoreboard freeze

Frozen batch-only standings compared with official final standings, including optimisation

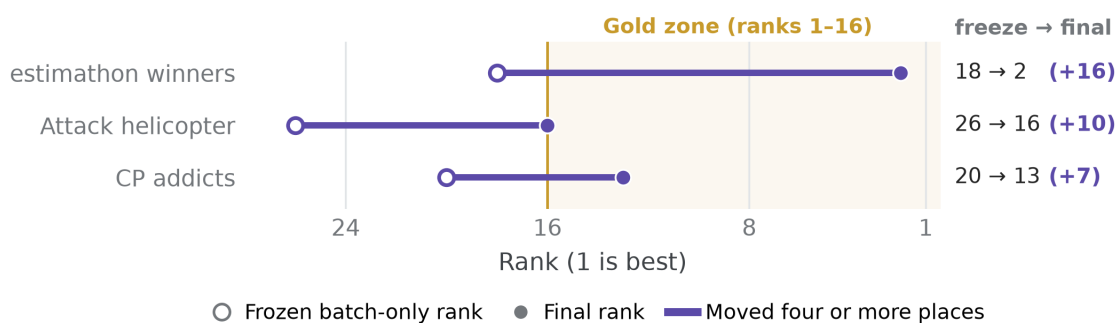


Figure 14: Rank changes from the frozen batch-only scoreboard to the official final standings for selected teams.

Figure 15: Final Contest leaderboard progression, sampled every five contest minutes, showing the top 16 teams with their batch and optimisation task scores. Compatible viewers provide pause, step and speed controls.

5 Post-Contest Survey Results

Shortly after each contest concluded, participants were invited to complete a feedback survey. We are grateful to the 111 respondents from the Online Contest (18.3% of the 606 participants in the official results) and the 47 respondents from the Final Contest (32.9% of the 143 attendees). Their feedback will help guide decisions for future editions of SIL. As this was the first SIL to include a physical event, their responses were especially valuable in helping us understand how the event can be improved in future. As respondents were self-selected, the percentages reported below refer to respondents rather than to all participants.

A few participants also shared informal feedback when approached after the Final Contest, but these responses may not reflect the views of all participants.

Aggregate survey responses may be found in Appendix [D](#).

5.1 Online Contest

The contest platform was well-received: For each of ease of use and smoothness, 82.9% of respondents gave a rating of 4 or 5 out of 5. 76.6% found the tasks interesting. 73.9% felt there were enough beginner-friendly tasks, 68.5% thought difficulty increased at a reasonable pace, and 71.2% found SIL useful as an introduction to competitive programming. However, only 42.3% rated their engagement with the optimisation task at 4 or 5, and only 56.8% felt that the task set was accessible overall.

5.2 Final Contest

5.2.1 Overall Reception

The Final Contest was received positively overall. 91.5% of respondents found the event well-organised, and 97.9% rated the food highly. Participants also responded positively to the contest content, with 78.7% finding the tasks interesting and 66.0% rating their engagement with the optimisation task at 4 or 5. These results suggest that both the event experience and the task set were appreciated by finalists.

Accessibility was somewhat weaker, with 57.4% of respondents feeling that there were enough beginner-friendly tasks. While the contest was intended to provide a greater challenge to this stronger field, the feedback suggests that there is room to provide more accessible entry points within the task set.

Only 48.9% of respondents rated the Final Contest highly as an introduction to computer science contests. This is perhaps unsurprising, given that the audience consisted of selected finalists who were likely to have greater prior exposure to competitive programming, and that the Final

Pre-contest resources used

Online survey, n = 79 answerers; respondents could select more than one resource.

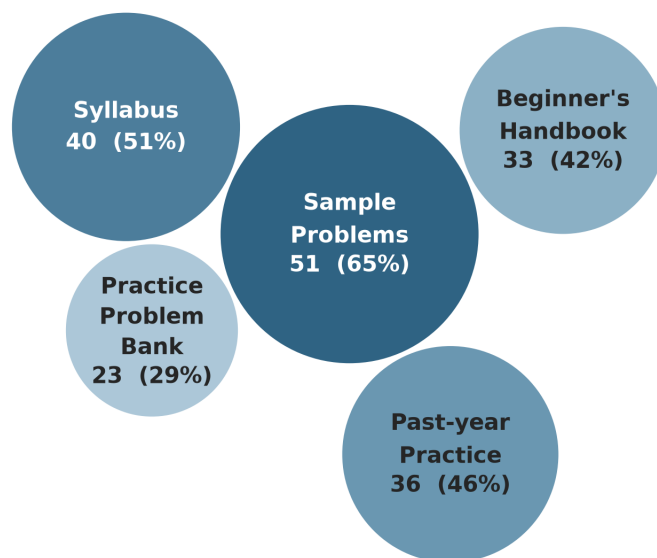


Figure 16: Resources used before SIL among the 79 Online Contest survey respondents who answered the multi-select question. Circle area is proportional to the number of respondents; choices were non-exclusive.

Contest tasks were more difficult. Nevertheless, the result remains useful when considering how the Final Contest can balance a challenging competition with an experience that remains educational and approachable, especially given the inclusion of eight wild-card entry teams.

Participants praised the on-site atmosphere, team experience and event organisation. Through informal feedback, nine of the ten people who discussed the Resolver described it positively. Four specifically highlighted the suspense and unpredictable standings. Several also valued SIL's exposure to the wider informatics community.

5.2.2 Issues Highlighted

Several teams noted that two laptops were insufficient, stating that sharing the laptops required participants to wait for their turn to use the laptop. We recognise that the unique team-based contest format is unfamiliar to participants. The format closely mirrors the ICPC, where three contestants share one laptop. Unfortunately, due to logistical constraints, we will likely be unable to increase the number of laptops available. We encourage teams to explore strategies to work efficiently as a team, such as writing code on paper when the laptops are unavailable.

Many contestants reported that the printed task statements were not easily separable, as the last page of a task might be printed on the same piece of paper as the first page of the subsequent

task. This issue arose from a miscommunication with printing services; our intention was to print task statements separately. We will ensure that this does not happen in future editions.

Three respondents raised concerns about the availability or suitability of vegetarian options during lunch. This was an oversight on our part, and we apologise to the affected participants.

One team told us that being called on stage individually as a lower-ranked team was a hurtful experience. We are deeply sorry for the hurt our decision caused, and we will commit to not calling lower-ranked teams on stage individually in future editions. Separately, we have reached out to the affected team directly to apologise and address their concerns.

There were two suggestions to allow participants to use their preferred name instead of their given name. We acknowledge that this is important to some participants, and are committed to implementing this change in future editions.

6 Evaluating SIL's Objectives

This section reviews the extent to which SIL 2026 has achieved its stated objectives. As described on the [SIL website](#), SIL aims to:

1. lower barriers to entry for students who are just starting their journey in computer science, by offering tasks that range from beginner-friendly to challenging;
2. serve as a stepping stone to the NOI;
3. emphasise teamwork and collaborative thinking;
4. welcome participants of all backgrounds, in particular those from traditionally underrepresented schools and genders.

Section 6.1 examines barriers to entry and participant engagement, Section 6.2 the progression of participants towards the NOI, Section 6.3 the extent to which teams competed collaboratively, and Section 6.4 representation and access to the Final Contest.

6.1 Lowering Barriers to Entry

Figure 17 suggests that engagement is strong across all groups for batch tasks at the start, with most teams continuing to improve during the first hour. This supports the solve rates in Figure 5, illustrating that Level 1 tasks were well-tuned to the majority of the teams, and that the contest is initially accessible for participants regardless of their experience.

After the first hour, the separation between the cohorts becomes clearer: all-new teams plateau earliest, while teams with some prior experience continue to improve for longer, and advanced

teams have both the highest scores and the greatest rate of improvement. The widening gap is consistent with the lower solve rates of the Level 2 tasks in Figure 5, suggesting that these tasks were substantially more effective at differentiating teams by experience.

True scores separated most after the first hour

142 eligible teams. Colour, line style and marker distinguish the cohort medians; the matching shaded band is that cohort's interquartile range. Scores are true batch scores through T+240; the shaded T+210–240 window is what the frozen public scoreboard did not show.

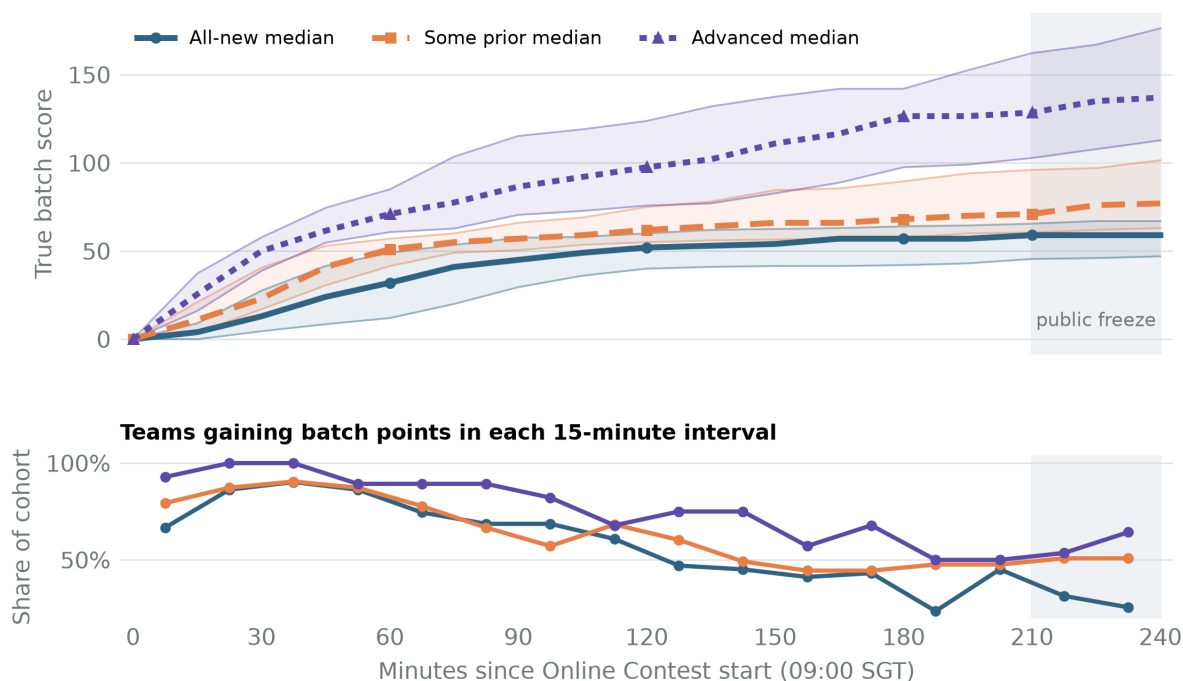


Figure 17: Online Contest batch-score progression by experience cohort, using true scores including the final 30 minutes hidden by the scoreboard freeze. In the upper panel, lines show cohort medians and bands show interquartile ranges.

Separately, the registration profile in Section 2 is consistent with this objective: more than half of participants reported no formal computing qualification, more than half reported a self-assigned skill level of 1 (Beginner), and the number of NOI newcomers grew from 2025. The survey results in Section 5 point in the same direction, with 73.9% of Online Contest respondents feeling that there were enough beginner-friendly tasks, 68.5% thinking that difficulty increased at a reasonable pace, and 71.2% finding SIL useful as an introduction to competitive programming.

However, only 56.8% of respondents felt that the task set was accessible overall, and several wild-card participants found the Final Contest tasks particularly difficult. Taken together, the evidence suggests that the Online Contest lowered barriers to entry for beginners, while the accessibility of the more challenging tasks, particularly at the Final Contest, remains an area for improvement; we will make the task set more approachable in 2027.

6.2 SIL as a Stepping Stone towards NOI

One of SIL's core aims is to serve as a stepping stone to the NOI by providing a more accessible competitive environment in which participants can build confidence and deepen their interest in informatics.

Among the 116 SIL returners who had no prior NOI experience in 2025, 41 (35.3%) went on to record some form of NOI participation in 2026 (Figure 18). In particular, 14 students achieved an NOI medal: 4 Bronze, 8 Silver, and 2 Gold. These outcomes are observed through the NOI experience field self-reported at SIL 2026 registration, and are therefore available only for the 206 returners who could be matched across both editions.

It should be noted that an absence of prior NOI experience does not imply an absence of prior competitive programming experience. Therefore, these results should not be interpreted as showing that SIL introduced these students to competitive programming, nor as evidence that SIL caused their subsequent NOI participation. Moreover, returners are a self-selected group who are likely to be more engaged than SIL 2025 participants as a whole, so the participation rate above should not be generalised to all SIL 2025 participants. Rather, these results show how SIL can be a stepping stone in participants' progression towards NOI, consistent with the role envisioned in SIL's mission. We encourage more SIL participants to continue their competitive programming journey by attending NOI.

14 NOI newcomers returned with a medal in 2026

Of 206 matched SIL returners, 116 had no NOI experience in 2025. Bars show their stored 2026 registration value; no matched fields were blank.

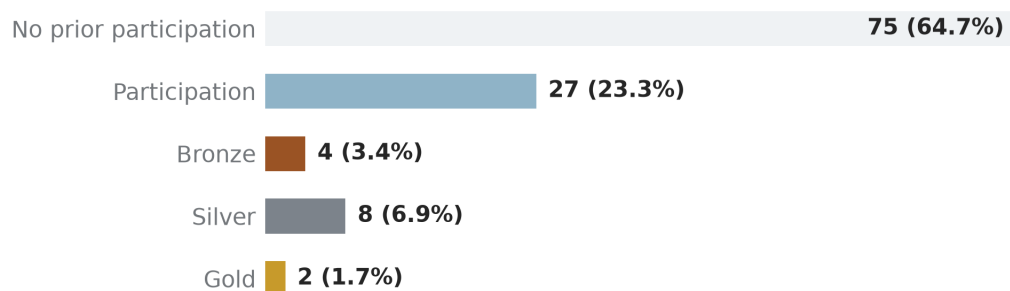


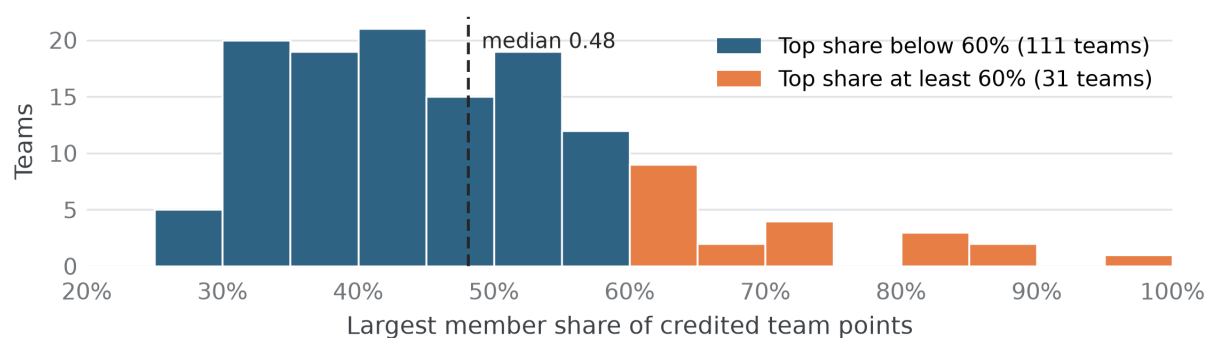
Figure 18: Self-reported NOI 2026 outcomes of the 116 SIL 2025 returners with no prior NOI experience.

6.3 SIL as a Team-based Competition

The extent to which teams collaborated cannot be measured directly. Instead, this section uses the distribution of scoring contributions within each team as a proxy: a balanced distribution cannot show that members worked together on the same tasks, but it can show whether a team's score reflected a team effort rather than the performance of a single member.

31 teams credited at least 60% of points to one member

Eligible Online teams (n = 142). Shares use points credited through each member's submissions; they do not measure all work contributed to the team.



52 of the 142 teams included at least one member with no credited points.

Figure 19: Distribution of the largest member share of credited points across the 142 eligible Online Contest teams.

Overall, contribution in the Online Contest remained reasonably well distributed across teams. A team is classified as carried when its top-scoring member accounted for at least 60% of the team's points. Only 31 of 142 teams were classified as being carried by a single member, while the median largest contributor accounted for 48% of their team's points. Hence, for most teams, no single member dominated the team's score and the contribution curves in Figure 20 for balanced teams show points being shared across multiple members.

However, contributions were not always evenly spread. 52 teams had at least one member who scored no points. In carried teams, the gap was much more obvious: contributions dropped sharply after the top member, with the other members often contributing very little.

Qualitative responses also highlighted the team and social experience, mentioned by 10 of the 77 Online Contest respondents and 11 of the 32 Final Contest respondents who left comments (Appendix D).

How credited points are distributed within teams

Median credited points relative to an even split, by member rank. Ranks 4 and 5 include 108 and 72 teams respectively.

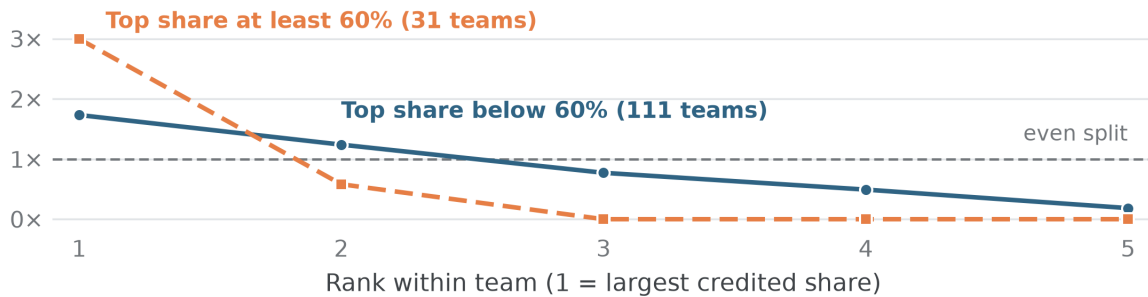


Figure 20: Median credited-point share by rank within the team across the 142 eligible Online Contest teams, shown as a multiple of an even split of team points. Ranks 4 and 5 apply only to teams with at least that many members.

6.4 Representation and Access to the Final Contest

A key objective of SIL is to promote computer science education among students from traditionally underrepresented schools and genders. In this respect, the results of SIL 2026 are mixed. Female, non-binary, other-gender, and undisclosed-gender participants formed 20.9% of registrations and 20.8% of the Online Contest field, but only 12.8% of registered finalists (19 of 148). Similarly, participants from schools with no NOI Gold medallist in the five most recent editions formed 54.0% of registrations and 48.8% of the Online Contest field, but only 20.3% of finalists (Table 1).

The difference is also visible at the team level: of the 51 teams in the official results whose members had no prior NOI experience, only 1 competed at the Final Contest, compared with 12 of the 63 teams with some prior experience and 19 of the 28 teams with advanced experience. As noted in Section 2.1, the proportion of female, non-binary, other-gender, and undisclosed-gender participants at registration also fell by 8.8 percentage points from 2025.

The wild-card qualification route was introduced to address this gap. In accordance with the rules, 8 of the 32 Final Contest slots were reserved for teams whose members were all female or non-binary, or where each member attended a school not represented among the teams that had already received an invitation. One team qualified because all of its members were female or non-binary, and seven because each of their members attended a school not otherwise represented.

Measure	Registrations	Registered Finalists	Direct and Replacement Qualifiers
Female, non-binary, other gender, or undisclosed	20.9%	12.8%	10.8%
From a school with no NOI Gold in five editions	54.0%	20.3%	7.2%

Table 1: Representation among the registered participants and the registered finalists. Registration shares are of the 757 registered participants; finalist shares are of the 148 registered finalists. The direct-and-replacement column covers the 111 of those finalists on the 24 teams that reached the Final Contest without a wild-card place.

Online rank plotted against Finals rank

Ranks are recomputed within the 32-team finalist cohort

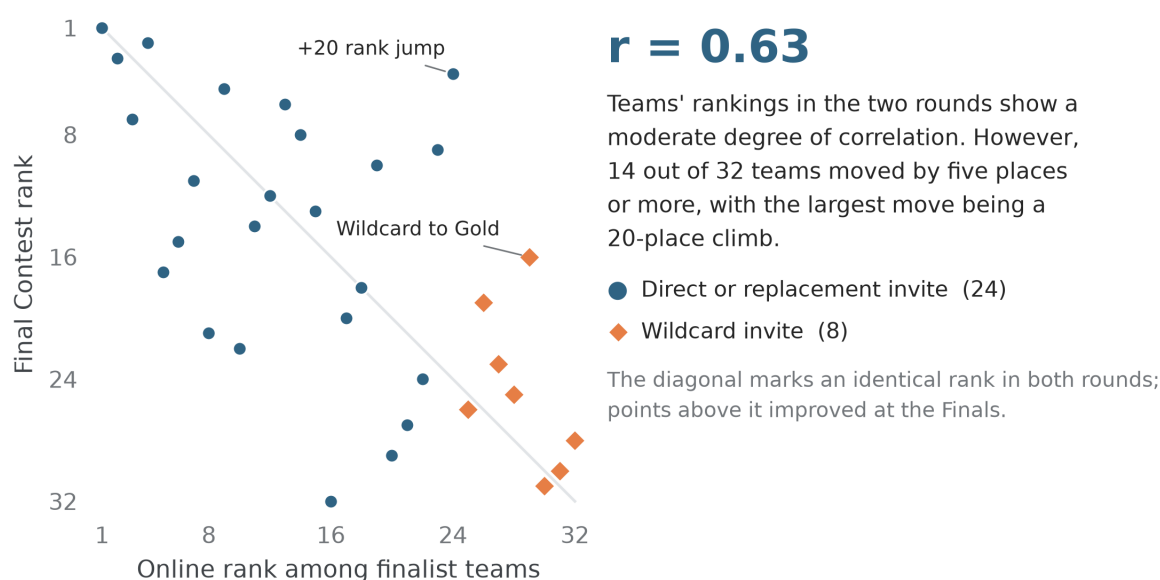


Figure 21: Online rank against Final Contest rank for the 32 teams that competed at the Final Contest. Online ranks are recomputed among finalists only; for example, a team ranked 50th in the full Online Contest may instead be ranked 30th within this cohort.

Although Online Contest rank is moderately positively correlated with Final Contest performance, there is still substantial variation between the two, with a rank correlation of 0.63; Online Contest performance is a useful indication of Final Contest performance, but it does not fully determine it.

Therefore, the wild-card entry system gives lower-ranked teams a chance to perform much better at the Final Contest than their Online Contest ranking might suggest, while also supporting SIL's goal of making the Final Contest more inclusive and accessible. We will maintain the wild-card qualification system in future editions.

Informal feedback from several participants in wild-carded teams was also positive. Some described the Final Contest as an eye-opening experience, particularly from being exposed to more experienced competitors, while others appreciated the qualification system, as it allowed more newcomer teams to experience and compete in the Final Contest. At the same time, multiple participants commented that they found the task set particularly difficult.

7 Conclusion

SIL 2026 was larger than its first edition and delivered its first on-site Final Contest. The two rounds were generally smoothly executed and well-received, but more can be done to improve the contest experience and fulfil our stated objectives. We are grateful for the enthusiasm and trust participants placed in our new format.

Our main priorities for 2027 follow from the findings of this report:

1. To broaden representation, we will maintain the wild-card qualification route, which brought eight teams from underrepresented groups to the Final Contest.
2. To improve the contest experience for all participants, we will make the task set more approachable.
3. To improve engagement with the optimisation task, we will introduce it more clearly.
4. To uphold the integrity of the contest, we will streamline and strengthen the post-contest verification review process.
5. To prevent easily avoidable errors in task statements or test data, we will tighten technical checks.
6. To improve the organisation of the Final Contest, we will act on the operational feedback in Section 5, including printing task statements separately and revising the closing ceremony.

We strongly encourage participants of SIL 2026 to continue learning about competitive programming. There are many free resources online, such as [CP Algorithms](#) and [USACO Guide](#). Participants can also consider competing on [Codeforces](#) and [AtCoder](#), or work towards Singapore's flagship competitive programming contest, the [National Olympiad in Informatics](#). Finally, we warmly invite graduating participants who would like to help organise SIL 2027 to fill in this [recruitment form](#).

A Participant Demographics

Category	Count	Percentage
Education Level		
Year 1 / Sec 1	20	2.64%
Year 2 / Sec 2	75	9.91%
Year 3 / Sec 3	142	18.76%
Year 4 / Sec 4 / Sec 5	103	13.61%
Year 5 / JC 1 / MI 1 / Poly 1 / ITE 1	226	29.85%
Year 6 / JC 2 / MI 2 / Poly 2 / ITE 2	172	22.72%
Poly 3 / MI 3 / ITE 3	19	2.51%
Gender		
Male	599	79.13%
Female	136	17.97%
Non-binary	2	0.26%
Other	5	0.66%
Prefer not to say	15	1.98%
NOI Experience		
None	485	64.07%
Certificate of Participation	158	20.87%
Bronze Medal	48	6.34%
Silver Medal	39	5.15%
Gold Medal	27	3.57%
Highest Computing Qualification		
None	391	51.65%
O-Level Computing	215	28.40%
Polytechnic	77	10.17%
H2 Computing	74	9.78%
Self-assigned Skill Level		
1 (Beginner)	407	53.76%
2	162	21.40%
3	132	17.44%
4	34	4.49%
5 (Veteran)	22	2.91%

Table 2: Breakdown of participant demographics and prior experience. $n = 757$ participants.

Registration-level Measure	2025	2026	Change
Registered teams	150	177	+27
Registered participants	620	757	+137
Represented schools [†]	56	70	+14
No previous NOI participation	69.35%	64.07%	−5.28 pp
	430	485	+55
No formal computing qualification	50.00%	51.65%	+1.65 pp
	310	391	+81
Self-assigned skill level 1	45.81%	53.76%	+7.95 pp
	284	407	+123
Female, non-binary, other gender, or undisclosed	29.68%	20.87%	−8.81 pp
	184	158	−26
International schools [†]	6	13	+7
Polytechnics [†]	4	5	+1
ITE institutions [†]	1	0	−1

Table 3: Compact demographic comparison using the published 2025 registration cohort and the 2026 registration cohort after removing the organiser test account. Percentage changes are percentage points. The 2025 counts are recovered from the published percentages of 620 registered participants, which is how the 2025 report reported them. [†]School-name normalisation differs between years, so institution counts are descriptive.

B Batch Task Statistics

Subtasks are scored independently; solving a later subtask does not imply solving earlier subtasks.

Subtask solve rates and wrong answers in the Online Contest

Online Contest: n = 142 eligible teams, grouped by published level. Only eligible teams are included.

Blue is the share of the cohort solving each subtask; grey is full solves per S1 solver; red is wrong answers — the average a full solver absorbed, and the share of all submissions rejected.

	S1	S2	S3	Full solve	S1 → full	Avg. WA (solved)	WA % (all)
1. Encluse Duck (L1)	98%	98%	98%	98%	100%	2.07	41%
2. Tag (L1)	99%	98%	97%	97%	99%	0.83	22%
3. Evil (L1)	98%	97%	96%	96%	99%	2.77	48%
4. Six Seven (L1)	99%	97%	96%	96%	97%	1.54	35%
5. Circles (L1)	96%	96%	96%	96%	100%	1.61	46%
6. Virus (L1)	94%	94%	94%	94%	100%	2.70	49%
7. ABBABBA (L1)	97%	97%	96%	96%	99%	2.32	44%
8. Block Game (L1)	97%	97%	97%	97%	100%	3.20	52%
9. Pringles Sort 2 (L2)	87%	73%	75%	72%	82%	1.93	74%
10. Duck Race (L2)	63%	54%	46%	46%	73%	0.97	78%
11. Heights (L2)	83%	57%	52%	52%	63%	0.85	40%
12. Mahjong 2 (L2)	55%	39%	37%	37%	67%	0.96	70%
13. Tokens (L2)	34%	21%	20%	20%	60%	2.59	90%
14. Eat Sleep Repeat (L2)	58%	37%	29%	29%	49%	1.54	77%
15. Bamboo 3 (L2)	65%	33%	31%	31%	47%	2.48	63%
16. Power Sum (L2)	47%	31%	23%	23%	49%	0.45	73%
17. Keys (L3)	44%	27%	20%	20%	45%	0.89	61%
18. Reels (L3)	14%	11%	9%	9%	65%	0.23	49%
19. Dictator (L3)	25%	17%	15%	15%	63%	0.91	48%
20. House Visit (L3)	10%	6%	6%	6%	64%	1.00	85%
21. Spinning Table (L3)	11%	9%	8%	8%	75%	0.67	66%
22. A Certain Scientific Committee (L3)	6%	5%	5%	5%	88%	0.71	61%
23. Bookshelf (L3)	10%	6%	6%	6%	64%	0.11	27%
24. Pawns (L3)	4%	1%	1%	1%	33%	0.00	91%
25. Turbines (L3)	5%	4%	4%	4%	71%	0.00	38%

Figure 22: Subtask and full-solve rates for each Online Contest batch task among all 142 eligible teams, grouped by Level.

Subtask solve rates and wrong answers in the Final Contest

Final Contest: $n = 32$ finalist teams, grouped by hidden level. Only eligible teams are included. Blue is the share of the cohort solving each subtask; grey is full solves per S1 solver; red is wrong answers — the average a full solver absorbed, and the share of all submissions rejected.

	S1	S2	S3	Full solve	S1 → full	Avg. WA (solved)	WA % (all)
1. 6 ate 8 (L2)	94%	81%	81%	81%	87%	1.00	36%
4. Reveille (L2)	100%	81%	81%	81%	81%	1.12	32%
6. Snowwabbits (L2)	94%	84%	78%	78%	83%	0.56	25%
2. Dog Walk (L3)	44%	50%	41%	38%	86%	0.33	38%
3. Gremlin Nob (L3)	47%	41%	16%	16%	33%	0.00	64%
9. Topsy Turvy (L3)	6%	9%	6%	6%	100%	0.00	63%
5. Slimes (L4)	0%	0%	0%	0%	n/a	n/a	100%
7. Spring Sunlight (L4)	28%	12%	6%	6%	22%	2.00	69%
8. StringStringString (L4)	12%	3%	3%	3%	25%	0.00	60%
10. Xiaoyang Ranking (L4)	6%	6%	6%	6%	100%	2.50	87%

Conversion is greyed where fewer than 5 teams solved Subtask 1, so the ratio rests on very few teams.

Figure 23: Subtask and full-solve rates for each Final Contest batch task among all 32 finalist teams, grouped by Level.

Table 4: Individual batch task submission statistics, Online Contest ($n = 142$ teams, after disqualifications). Avg. WA (Solved): Average number of Wrong Answers for teams that solved the subtask. WA % (All): Percentage of total submissions that were WAs.

Task / Subtask	Avg. WA (Solved)	WA % (All)	Solves	Solve %
Task 1: Enclose Duck (Full Solves: 139 97.89%)				
Subtask 1	0.91	47.74%	139	97.89%
Subtask 2	0.27	21.02%	139	97.89%
Subtask 3	0.89	47.15%	139	97.89%
Task 2: Tag (Full Solves: 138 97.18%)				
Subtask 1	0.14	12.50%	140	98.59%
Subtask 2	0.17	14.20%	139	97.89%
Subtask 3	0.53	36.11%	138	97.18%
Task 3: Evil (Full Solves: 137 96.48%)				
Subtask 1	1.17	54.13%	139	97.89%
Subtask 2	0.79	45.02%	138	97.18%
Subtask 3	0.82	44.98%	137	96.48%

Continued on next page...

Table 4 – Continued from previous page

Task / Subtask	Avg. WA (Solved)	WA % (All)	Solves	Solve %
Task 4: Six Seven (Full Solves: 136 95.77%)				
Subtask 1	0.71	41.91%	140	98.59%
Subtask 2	0.38	28.87%	138	97.18%
Subtask 3	0.48	33.98%	136	95.77%
Task 5: Circles (Full Solves: 137 96.48%)				
Subtask 1	0.55	49.82%	137	96.48%
Subtask 2	0.74	49.45%	137	96.48%
Subtask 3	0.31	35.07%	137	96.48%
Task 6: Virus (Full Solves: 134 94.37%)				
Subtask 1	1.72	65.46%	134	94.37%
Subtask 2	0.44	30.57%	134	94.37%
Subtask 3	0.54	35.27%	134	94.37%
Task 7: ABBABBA (Full Solves: 137 96.48%)				
Subtask 1	0.56	35.81%	138	97.18%
Subtask 2	0.84	45.67%	138	97.18%
Subtask 3	0.94	49.45%	137	96.48%
Task 8: Block Game (Full Solves: 138 97.18%)				
Subtask 1	0.20	16.87%	138	97.18%
Subtask 2	0.33	24.59%	138	97.18%
Subtask 3	2.67	72.73%	138	97.18%
Task 9: Pringles Sort 2 (Full Solves: 102 71.83%)				
Subtask 1	1.49	66.49%	124	87.32%
Subtask 2	1.36	84.94%	103	72.54%
Subtask 3	0.10	55.46%	106	74.65%
Task 10: Duck Race (Full Solves: 66 46.48%)				
Subtask 1	2.87	86.40%	90	63.38%
Subtask 2	0.34	64.15%	76	53.52%
Subtask 3	0.12	59.76%	66	46.48%
Continued on next page...				

Table 4 – Continued from previous page

Task / Subtask	Avg. WA (Solved)	WA % (All)	Solves	Solve %
Task 11: Heights (Full Solves: 74 52.11%)				
Subtask 1	0.22	26.25%	118	83.10%
Subtask 2	0.63	55.74%	81	57.04%
Subtask 3	0.16	32.73%	74	52.11%
Task 12: Mahjong 2 (Full Solves: 52 36.62%)				
Subtask 1	1.09	80.69%	78	54.93%
Subtask 2	0.27	60.43%	55	38.73%
Subtask 3	0.08	24.64%	52	36.62%
Task 13: Tokens (Full Solves: 29 20.42%)				
Subtask 1	3.40	94.50%	48	33.80%
Subtask 2	0.23	76.00%	30	21.13%
Subtask 3	0.10	54.69%	29	20.42%
Task 14: Eat Sleep Repeat (Full Solves: 41 28.87%)				
Subtask 1	0.73	80.29%	83	58.45%
Subtask 2	0.73	78.15%	52	36.62%
Subtask 3	0.46	58.16%	41	28.87%
Task 15: Bamboo 3 (Full Solves: 44 30.99%)				
Subtask 1	1.10	55.71%	93	65.49%
Subtask 2	1.30	75.26%	47	33.10%
Subtask 3	0.75	56.44%	44	30.99%
Task 16: Power Sum (Full Solves: 33 23.24%)				
Subtask 1	0.39	79.19%	67	47.18%
Subtask 2	0.09	58.10%	44	30.99%
Subtask 3	0.03	70.27%	33	23.24%
Task 17: Keys (Full Solves: 28 19.72%)				
Subtask 1	1.21	54.74%	62	43.66%
Subtask 2	0.21	73.47%	39	27.46%
Subtask 3	0.00	34.88%	28	19.72%
Continued on next page...				

Table 4 – Continued from previous page

Task / Subtask	Avg. WA (Solved)	WA % (All)	Solves	Solve %
Task 18: Reels (Full Solves: 13 9.15%)				
Subtask 1	0.05	65.52%	20	14.08%
Subtask 2	0.13	31.82%	15	10.56%
Subtask 3	0.00	13.33%	13	9.15%
Task 19: Dictator (Full Solves: 22 15.49%)				
Subtask 1	0.57	52.05%	35	24.65%
Subtask 2	0.25	57.14%	24	16.90%
Subtask 3	0.05	15.38%	22	15.49%
Task 20: House Visit (Full Solves: 9 6.34%)				
Subtask 1	0.57	91.62%	14	9.86%
Subtask 2	0.67	68.97%	9	6.34%
Subtask 3	0.11	30.77%	9	6.34%
Task 21: Spinning Table (Full Solves: 12 8.45%)				
Subtask 1	0.44	83.16%	16	11.27%
Subtask 2	0.08	13.33%	13	9.15%
Subtask 3	0.00	0.00%	12	8.45%
Task 22: A Certain Scientific Committee (Full Solves: 7 4.93%)				
Subtask 1	0.25	78.95%	8	5.63%
Subtask 2	0.43	36.36%	7	4.93%
Subtask 3	0.00	0.00%	7	4.93%
Task 23: Bookshelf (Full Solves: 9 6.34%)				
Subtask 1	0.07	30.00%	14	9.86%
Subtask 2	0.00	40.00%	9	6.34%
Subtask 3	0.00	0.00%	9	6.34%
Task 24: Pawns (Full Solves: 2 1.41%)				
Subtask 1	0.50	94.06%	6	4.23%
Subtask 2	0.00	50.00%	2	1.41%
Subtask 3	0.00	0.00%	2	1.41%
Continued on next page...				

Table 4 – Continued from previous page

Task / Subtask	Avg. WA (Solved)	WA % (All)	Solves	Solve %
Task 25: Turbines (Full Solves: 5 3.52%)				
Subtask 1	0.29	22.22%	7	4.93%
Subtask 2	0.00	45.45%	6	4.23%
Subtask 3	0.00	44.44%	5	3.52%

Table 5: Individual batch task submission statistics, Final Contest ($n = 32$ finalist teams). Avg. WA (Solved): Average number of Wrong Answers for teams that solved the subtask. WA % (All): Percentage of total submissions that were WAs.

Task / Subtask	Avg. WA (Solved)	WA % (All)	Solves	Solve %
Task 1: 6 ate 8 (Full Solves: 26 81.25%)				
Subtask 1	0.03	18.92%	30	93.75%
Subtask 2	0.88	54.41%	26	81.25%
Subtask 3	0.12	15.62%	26	81.25%
Task 2: Dog Walk (Full Solves: 12 37.50%)				
Subtask 1	0.14	61.11%	14	43.75%
Subtask 2	0.12	20.00%	16	50.00%
Subtask 3	0.00	0.00%	13	40.62%
Task 3: Gremlin Nob (Full Solves: 5 15.62%)				
Subtask 1	0.13	71.15%	15	46.88%
Subtask 2	0.15	45.83%	13	40.62%
Subtask 3	0.00	68.75%	5	15.62%
Task 4: Reveille (Full Solves: 26 81.25%)				
Subtask 1	0.25	20.00%	32	100.00%
Subtask 2	0.88	53.57%	26	81.25%
Subtask 3	0.00	3.70%	26	81.25%
Task 5: Slimes (Full Solves: 0 0.00%)				
Subtask 1	N/A	100.00%	0	0.00%
Subtask 2	N/A	100.00%	0	0.00%
Subtask 3	N/A	100.00%	0	0.00%

Continued on next page...

Table 5 – Continued from previous page

Task / Subtask	Avg. WA (Solved)	WA % (All)	Solves	Solve %
Task 6: Snowwabbits (Full Solves: 25 78.12%)				
Subtask 1	0.37	30.23%	30	93.75%
Subtask 2	0.26	22.86%	27	84.38%
Subtask 3	0.12	21.88%	25	78.12%
Task 7: Spring Sunlight (Full Solves: 2 6.25%)				
Subtask 1	0.56	64.00%	9	28.12%
Subtask 2	0.25	77.78%	4	12.50%
Subtask 3	0.00	60.00%	2	6.25%
Task 8: StringStringString (Full Solves: 1 3.12%)				
Subtask 1	0.00	60.00%	4	12.50%
Subtask 2	0.00	50.00%	1	3.12%
Subtask 3	0.00	66.67%	1	3.12%
Task 9: Topsy Turvy (Full Solves: 2 6.25%)				
Subtask 1	0.00	66.67%	2	6.25%
Subtask 2	0.00	66.67%	3	9.38%
Subtask 3	0.00	50.00%	2	6.25%
Task 10: Xiaoyang Ranking (Full Solves: 2 6.25%)				
Subtask 1	1.50	91.30%	2	6.25%
Subtask 2	1.00	66.67%	2	6.25%
Subtask 3	0.00	87.50%	2	6.25%

C Optimisation Statistics

Round	Subtask	Teams Attempting	Field	Median Raw	Best Raw
Online Contest	1	74	142	48.5	49
Online Contest	2	68	142	39,518	54,062
Online Contest	3	55	142	37,305	94,254
Final Contest	1	32	32	477	480
Final Contest	2	32	32	10,937.5	12,420
Final Contest	3	31	32	781,143	1,360,701

Table 6: Optimisation attempts and raw-score summaries. Field is the number of eligible teams in that round; a team counts as attempting a subtask if it recorded any raw score on it.

Online Contest: raw optimisation scores

Ticks below the baseline are individual teams; the dashed rule marks the median.

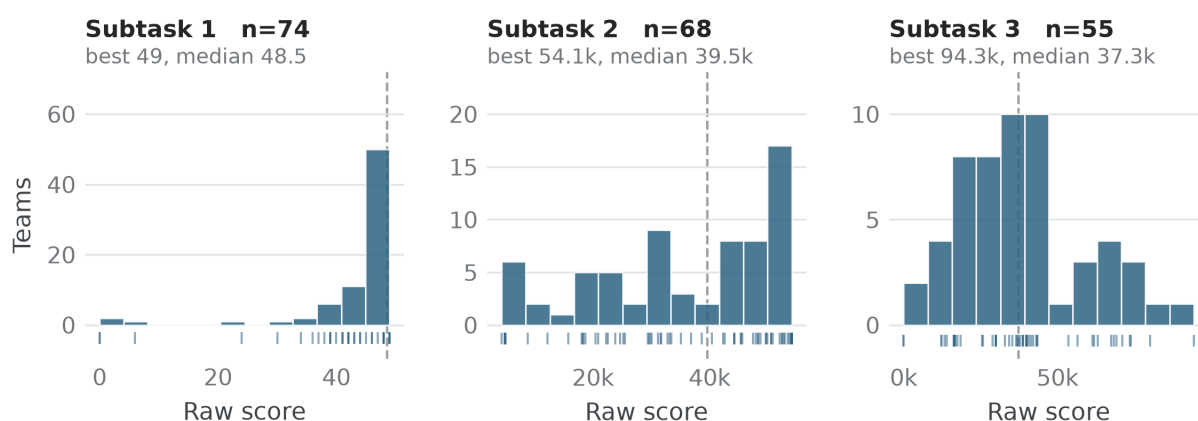


Figure 24: Raw optimisation scores for each Online Contest subtask. Subtask 1 piles up against its ceiling; Subtasks 2 and 3 spread teams out. Each panel has its own scale, and the dashed vertical line marks the median.

Final Contest: raw optimisation scores

Ticks below the baseline are individual teams; the dashed rule marks the median.

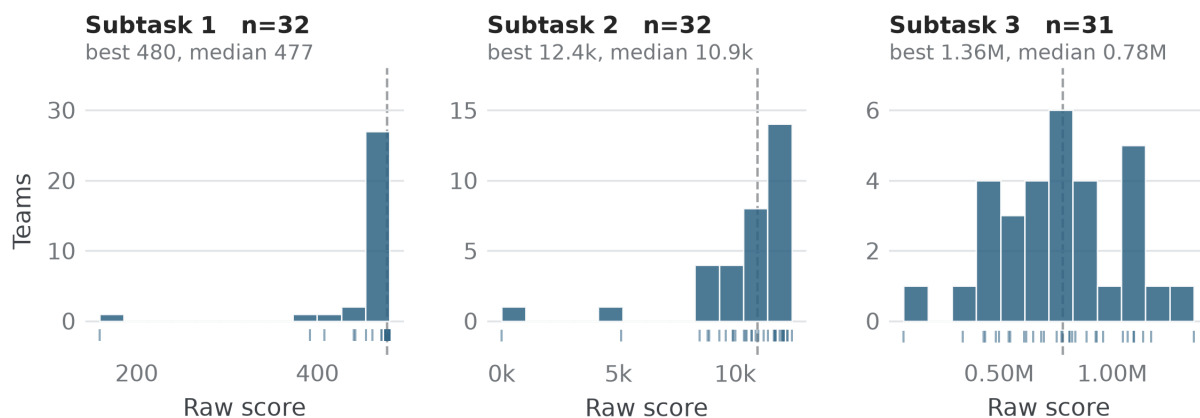


Figure 25: Raw optimisation scores for each Final Contest subtask. Subtask 1 again concentrates near the best construction, while Subtask 3 spreads widely. Each panel has its own scale, and the dashed vertical line marks the median.

D Survey and Qualitative Feedback

Question	1 (Low)	2	3	4	5 (High)
Task Interest	0.9%	3.6%	18.9%	49.5%	27.0%
Accessibility (Difficulty)	1.8%	4.5%	36.9%	42.3%	14.4%
Optimisation Engagement	2.7%	3.6%	51.4%	23.4%	18.9%
Platform Ease of Use	0.9%	0.9%	15.3%	32.4%	50.5%
User Experience (Smoothness)	2.7%	1.8%	12.6%	34.2%	48.6%
CS Contest Introduction Value	2.7%	7.2%	18.9%	39.6%	31.5%
CS/Programming Learning Value	5.4%	6.3%	19.8%	36.0%	32.4%

Table 7: Percentage distribution of participant ratings across contest metrics, Online Contest post-contest survey ($n = 111$ respondents).

Question	Response Option	Percentage
Reasonable Difficulty Pace	Yes / No	68.5% / 31.5%
Enough Beginner Tasks	Yes / No	73.9% / 26.1%
Contest Duration (4 Hours)	Just right	67.6%
	Too short	28.8%
	Too long	3.6%

Table 8: Binary and categorical responses regarding contest structure, Online Contest post-contest survey ($n = 111$ respondents).

Question	1 (Low)	2	3	4	5 (High)
Event Organisation	4.3%	2.1%	2.1%	34.0%	57.4%
Food Provided	0.0%	0.0%	2.1%	31.9%	66.0%
Task Interest	2.1%	2.1%	17.0%	53.2%	25.5%
Accessibility (Difficulty)	4.3%	2.1%	34.0%	42.6%	17.0%
Optimisation Engagement	2.1%	8.5%	23.4%	25.5%	40.4%
Platform Ease of Use	4.3%	2.1%	4.3%	48.9%	40.4%
User Experience (Smoothness)	2.1%	6.4%	4.3%	42.6%	44.7%
CS Contest Introduction Value	2.1%	17.0%	31.9%	27.7%	21.3%
CS/Programming Learning Value	4.3%	10.6%	29.8%	31.9%	23.4%

Table 9: Percentage distribution of participant ratings across contest metrics, Final Contest post-contest survey ($n = 47$ respondents).

Question	Response Option	Percentage
Reasonable Difficulty Pace	Yes / No	61.7% / 38.3%
Enough Beginner Tasks	Yes / No	57.4% / 42.6%
Contest Duration (4 Hours)	Just right	74.5%
	Too short	25.5%
	Too long	0.0%

Table 10: Binary and categorical responses regarding contest structure, Final Contest post-contest survey ($n = 47$ respondents).

Theme	Respondents	Percentage
Online Contest survey		
Platform or logistics	45	58.4%
Difficulty or pacing	28	36.4%
Different from other contests	24	31.2%
Optimisation task	20	26.0%
Enjoyment or motivation to continue	15	19.5%
Team or social experience	10	13.0%
New to informatics or contests	3	3.9%
Learned something specific	1	1.3%
Meeting the wider community	0	0.0%
Final Contest survey		
Platform or logistics	23	71.9%
Team or social experience	11	34.4%
Different from other contests	9	28.1%
Optimisation task	9	28.1%
Enjoyment or motivation to continue	5	15.6%
Difficulty or pacing	5	15.6%
New to informatics or contests	0	0.0%
Learned something specific	0	0.0%
Meeting the wider community	0	0.0%

Table 11: Themes in the optional free-text survey answers, coded after contextual review. Each respondent is counted once per theme; a respondent may raise several. Percentages are of the respondents who left free text ($n = 77$ for the Online Contest survey and $n = 32$ for the Final Contest survey).

E Task Topic Classification

Task	Name	Level	Primary Topic	Secondary Topic
9	Pringles Sort 2	2	Ad-hoc	–
10	Duck Race	2	Greedy	Mathematics
11	Heights	2	Ad-hoc	Sortings
12	Mahjong 2	2	Greedy	Sortings
13	Tokens	2	Greedy	Dynamic Programming
14	Eat Sleep Repeat	2	Greedy	Sortings
15	Bamboo 3	2	Ad-hoc	Mathematics
16	Power Sum	2	Brute Force	Data Structures
17	Keys	3	Binary Search	Data Structures
18	Reels	3	Greedy	Data Structures
19	Dictator	3	Dynamic Programming	Binary Search
20	House Visit	3	Classic Techniques	Sortings
21	Spinning Table	3	Mathematics	Classic Techniques
22	A Certain Scientific Committee	3	Sortings	Classic Techniques
23	Bookshelf	3	Ad-hoc	Games & Constructive
24	Pawns	3	Games & Constructive	Greedy
25	Turbines	3	Graph Theory	Greedy

Table 12: Online Contest task topic classification. Level 1 tasks were intentionally excluded.

Task	Name	Level	Primary Topic	Secondary Topic
1	6 ate 8	2	Greedy	Data Structures
2	Dog Walk	3	Graph Theory	Greedy
3	Gremlin Nob	3	Classic Techniques	–
4	Reveille	2	Greedy	–
5	Slimes	4	Data Structures	Brute Force
6	Snowwabbits	2	Sortings	–
7	Spring Sunlight	4	Classic Techniques	Binary Search
8	StringStringString	4	Graph Theory	Classic Techniques
9	Topsy Turvy	3	Greedy	Binary Search
10	Xiaoyang Ranking	4	Greedy	Binary Search

Table 13: Final Contest task topic classification.