

Debating & Judging Manual

(updated by Sofia WUDC 2026 CAP)

A Note about the Authorship of This Manual

SU Debaters: We mainly do YODL debating which takes a more "good reasons/judgment/argument" approach rather than the rules approach of the Worlds Manual.

The World Universities Debating Championships ('WUDC' or 'Worlds') Debating and Judging Manual was initially compiled in advance of the 35th World Championships in Malaysia. This Manual is the product of the time, work, wisdom and effort of many CAPs and debating intellectuals and academics.

The Sofia WUDC 2026 CAP (Marta Vasić, Matt Caito, Anna Shreder, Chris Mentis Cravaris, David Africa, Dennis Su, Joy Hadome, Matthew Toomey, Ruth Selorme Acolaste, and Sajid Asbat Khandaker) have worked together to update and clarify sections of this document.

Abridged List of Substantial Prior Contributors

- **The Panama WUDC 2025 CAP:** Juanita Hincapié Restrepo, Erasmus Mawuli Segbefia, Jane Mentzinger, Kat Cheng, Marta Vasić, Pranav Kagalkar, Shudipto R. Ahmed
- **The Madrid WUDC 2023 CAP:** Jessica Musulin, Sourodip Paul, Klaudia Maciejewska, Nar Sher May, Njuguna Macharia, Ruth Silcoff
- **The Belgrade WUDC 2022 CAP:** Brent Schmidt, Enting Lee, Hadar Goldberg, Juanita Hincapié Restrepo, Milos Marjanovic, Noluthando Honono, Robert Barrie, Yarn Shih
- **The Korea WUDC 2021 CAP:** Bobbi Leet, Boemo Phirinyane, Connor O'Brien, Dan Lahav, Milos Marjanovic, Mubarrat Wassey, Sebastian Dasso, Sooyoung Park, Tejal Patwardhan, Teck Wei Tan
- **The Thailand WUDC 2020 CAP:** Ashish Kumar, Jasmine Ho, Michael Dunn-Goekjian, Ayal Hayut-Man, Archie Hall, Mayu Arimoto, Lucia Arce, Julio Meyer, Evan Lynyak, Cliff Simataa, Jessica Musulin
- **The Cape Town WUDC 2019 CAP:** Ameera Natasha Moore, Fanele Mashwama, Dan Lahav, Elisha Kunene, Enting Lee, Raffy Marshall and Steph White
- **The Mexico WUDC 2018 CAP:** Michael Dunn-Goekjian, Mifzal Mohammed, Emma Johnstone, Sella Nevo, Wasifa Noshin, Steven Penner, Paul Smith, Evan Lynyak and Nicholas Ferezin
- **The Dutch WUDC 2017 CAP:** Syed Saddiq, Karin Merckens, Hyewon Rho, Michael Dunn-Goekjian, Omer Nevo, Jodie O'Neill, Veenu Goswami, Yashodhan Nair and Ingrid Rodriguez
- **The Thessaloniki WUDC 2016 CAP:** Manos Moschopoulos, Arinah Najwa, Chris Bisset, John McKee, Josh Zoffer, Sarah Balakrishnan and Tasneem Elias
- **The Malaysia Worlds 2015 Adjudication Team:** Shafiq Bazari, Jonathan Leader Maynard, Engin Arkan, Brett Frazer, Madeline Schultz, Sebastian Templeton, Danique van

Kopenhagen

- **Past WUDC and EUDC Chief Adjudicators:** Michael Baer, Sam Block, Doug Cochran, Lucinda David, Harish Natarajan, Sharmila Parmanand, Hadar Goldberg, Milos Marjanovic, Geneva Roy, Jack Stanley, Velina Andonova
- **The Athens EUDC 2019 CAP:** Dan Lahav, Sharmila Parmanand, Benji Kalman, Brian Wong, Cliodhna Ni Cheileachair, and Milos Marjanovic for their creation of the judge feedback scale
- **The Warsaw EUDC 2016 CAP:** Emilia Carlqvist, Harish Natarajan, Adam Hawksbee, Helena Ivanov, Radu Cotarcea, and Yael Bezalel for their updates to the speaker scale.
- **The Manchester EUDC 2013 CAP:** Alex Worsnip, Andrew Tuffin, Dessislava Kirova, Filip Dobranić, Omer Nevo, Joe Roussos, Jonathan Leader Maynard, Sam Block, and Shengwu Li for their original work on the WUDC Speaker Scale
- **The Cork WUDC 1996/Stellenbosch WUDC 1997 Rules Committee:** Ray D'Cruz, Tommy Tonner, Bob Dalrymple, Brent Patterson, Sacha Judd, John Long, Meg O'Sullivan, and Omar Salahuddin for their original work in formulating a standard set of rules for WUDC.

Table of Contents

Introduction.....	5
1. Core Rules of BP Debating.....	5
1.1 BP Format.....	5
1.2 Length of Speeches.....	6
1.3 Points of Information.....	6
Considering POIs in Judging Engagement.....	7
Cutting off a POI.....	8
Barracking/Badgering.....	8
Points of Clarification.....	8
1.4 Before the Debate.....	9
The Motion.....	9
Information, Context, or Definitions Accompanying Motions.....	9
Preparation Time.....	9
Pronoun Introductions.....	10
1.5 Iron-personing.....	11
1.6 Breaches of Order.....	11
Calling Order.....	12
Stopping the Clock.....	12
1.7 Tournament Structure.....	12
2. Debating and Judging at WUDC.....	14
2.1 Winning a Debate.....	14
2.2 The ‘Ordinary Intelligent Voter’.....	14
Facts, Knowledge, and Special Language.....	14
Dispositions.....	15
Judging as the Ordinary Intelligent Voter.....	15
2.3 Persuasiveness.....	16
Analysis.....	17
Style.....	18
2.4 Contradictions.....	19
What Is a Contradiction?.....	19
Contradictions within the Same Speech or within the Same Team.....	20
Contradictions between Teams on the Same Bench.....	20
How Teams Should Deal with Contradictions from the Other Side.....	20
2.5 Rebuttal, Engagement, and Comparisons.....	21
2.6 Burdens.....	22
Weighing Competing Frameworks.....	23
2.7 Motion Types.....	24
Policy Motions.....	24
Analysis Motions.....	25
Actor Motions.....	28
Alternative Motions.....	30
2.8 Role Fulfilment.....	30

2.9 Definitions and Models.....	31
Squirrelling.....	32
Vague definitions.....	33
Challenging the Definition.....	34
2.10 Opposing the Debate.....	36
Counter-propping.....	36
2.11 Member Speeches - Extending the Debate.....	39
Knifing.....	40
2.12 Whip Speeches.....	41
2.13 Equity.....	41
3. Additional Notes for Judges.....	43
3.1 Deciding the Results.....	43
3.2 Judging Panels.....	44
Voting on Pairwise Comparisons.....	45
Cyclic Ties.....	45
Example A - Three-person Panel.....	45
The Ranked Pairs Method.....	46
Ranked Pairs Explained.....	46
Example B - Five-person Panel.....	46
Trainee Judges.....	49
3.3 Managing the Discussion.....	49
3.4 Filling in the Ballot.....	50
3.5 Announcing the Result.....	51
3.6 Feedback on Adjudicators.....	52
Appendix A: The WUDC Speaker Scale.....	52
Appendix B: Chair Feedback Scale.....	56
Appendix C: Panellist and Trainee Feedback Scale.....	60

Introduction

This manual is divided into three chapters. Chapter One explains the fundamental format and operation of debates in the British Parliamentary ('BP') format used at Worlds. Chapter Two, explains how judges should evaluate debaters and, consequently, how debaters ought to debate. Chapter Three offers some additional notes for judges, covering issues like how the deliberation process works, assigning speaker points, and giving feedback.

1. Core Rules of BP Debating

1.1 BP Format

Each debate will contain four teams, each team consisting of two speakers.

There are two teams on each side of the debate. On one side are Opening Government ('OG') and Closing Government ('CG'), on the other side are Opening Opposition ('OO') and Closing Opposition ('CO').

The two sides of the debate are sometimes called 'benches' - as in, 'the Government bench' and 'the Opposition bench'. The first two teams in the debate (OG and OO) are sometimes collectively called the 'opening half', whilst the third and fourth teams in the debate (CG and CO) are sometimes collectively called the 'closing half'.

	Government Bench	Opposition Bench
Opening Half	OG • Prime Minister ('PM') • Deputy Prime Minister ('DPM')	OO • Leader of the Opposition ('LO') • Deputy Leader of the Opposition ('DLO')
Closing Half	CG • Member of Government ('MG') • Government Whip ('GW')	CO • Member of Opposition ('MO') • Opposition Whip ('OW')

In the order specified below, speakers from the four teams give their speeches, with each speaker giving one speech:

1. First speaker (the 'Prime Minister') from the OG team,

2. First speaker (the 'Leader of the Opposition') from the OO team,
3. Second speaker (the 'Deputy Prime Minister') from the OG team,
4. Second speaker (the 'Deputy Leader of Opposition') from the OO team,
5. First speaker (the 'Member of Government') from the CG team,
6. First speaker (the 'Member of Opposition') from the CO team,
7. Second speaker (the 'Government Whip') from the CG team,
8. Second speaker (the 'Opposition Whip') of the CO team.

The debate is presided over by a 'Chair', a designated individual who oversees the proceedings of the debate, calling on speakers to speak and enforcing the rules. Each debate will also usually have a timekeeper, who could be the Chair, another judge, or another individual entirely, who times speakers' speeches.

1.2 Length of Speeches

Speeches last for 7 minutes. Time signals (usually a bang on the table, ring of a bell, or clap of the hands) will be given by the timekeeper to indicate when 1 minute, 6 minutes and 7 minutes (often indicated by a double clap/bang) have elapsed. Though speakers should ideally finish their speech by 7 minutes, they may legitimately continue to speak in order to finish their sentence or wrap up a conclusion. As a general rule, this shouldn't take more than a further 15 seconds. Beyond 7 minutes and 15 seconds, judges are no longer permitted to take **anything** the speaker says into account. The Chair or timekeeper of the debate should bang the table or clap three times at 10 second intervals after 7 minutes 15 seconds to remind the speaker that they are now well beyond their time limit. If the speaker continues speaking past 7 minutes 30 (which should never happen), the Chair of the debate should 'call order', and instruct the speaker to sit down.

Speakers should start their speech as soon as they are called on by the chair of the debate, unless in reasonable circumstances as approved by the chair. Speakers may take reasonable time to organise their notes and start their timer. Chairs should ensure the debate proceeds in a timely manner.

1.3 Points of Information

A POI is a formalised interjection from any speaker on the opposite side of the bench to the speaker who has the floor. It is up to the speaker who has the floor to decide which POIs to accept (i.e., allow to be made) or reject (i.e., not allowed to be made).

The first and last minute of each speech is known as 'protected time', during which no POIs may be offered to the speaker who is making their speech. During the intervening 5 minutes (i.e., between 1 minute and 6 minutes) POIs may be offered. **Each speaker must take at least one POI or they shall be penalised for failure to engage.**

A POI may last up to 15 seconds. It can take the form of a comment or a question to the speaker who has the floor. To offer a POI, a speaker should say 'point of information,' 'on that point' or

‘point’. They should not offer ‘coded POIs’ by uttering anything which reveals the content of the POI before it has been accepted (by saying, for example ‘on the law’ or ‘not at all!’). If the POI offered is refused, the speaker who offered it should sit down immediately.

POIs may not be offered after the 6 minute mark in a speaker’s speech, and at 6 minutes all speakers currently standing (to indicate that they have offered a POI) should sit down. It is acceptable for a POI which was offered and accepted before the 6 minute mark to continue to be made past the 6 minute mark - it should continue until the POI is concluded, the 15 second time allotment has passed or the POI is cut off. It is also acceptable for a POI offered before 6 minutes to be accepted by a speaker dead on the 6 minute mark and then be made. Once all speakers are sitting after the 6 minute mark, no more POIs may be offered or accepted.

Sometimes speakers may express a ‘POI preference’ such as:

- demanding certain speakers or teams stop offering POIs
- saying they will only take a POI from a specific team (e.g., PM says they will only take a POI from CO)
- asking speakers or teams to only ask POIs at a specific time (e.g., after the 5th minute)
- asking speakers or teams to ask non-verbal POIs (e.g., only ask by raising their hand).

However, debaters are **not obligated to follow these preferences**. Additionally, **judges should not enforce these preferences, and their judgement of the debate should not be altered by the expressed preferences**. All debaters have the right, throughout the times the rules allow in the debate, to offer POIs to speakers from the other side. Similarly, a speaker calling for a POI to be offered does not create any special obligation for a team or speaker to offer a point.

Considering POIs in Judging Engagement

POIs are an important component in debate rounds. It is the responsibility of judges to track and evaluate POI engagement during the round, which includes but is not limited to: whether or not a speaker was offered POIs, whether or not a speaker accepted a POI, the quality of the POI asked as well as the quality of the POI response. If a speaker has not accepted a POI, judges must remind the room to accept POIs after the speaker has finished speaking.

When evaluating speakers that have not taken POIs (assuming sufficient POIs were offered), judges must treat a failure to take a POI as indicative of a reduced level of engagement and evaluate this as reducing the persuasiveness of that speaker’s contribution. For instance, judges should lower speaker scores for the speaker that did not accept POIs to reflect their reduced level of engagement, adjust the margin of victory for teams, or flip close calls between teams. This **does not** mean that a team will take an automatic fourth for failing to take a POI, **nor does it mean** that they cannot win the debate!

If a speaker was offered no POIs, or was only offered one or two POIs at the start of their speech and had no opportunities to take POIs towards the later half of their speech, they will **not be penalised for a lack of engagement** (after all, it is difficult to engage when there is nothing to engage with). A speaker in such circumstances may explicitly ask for a POI, and doing so will demonstrate a willingness to engage with arguments even if no POI is subsequently offered.

In general, judges should evaluate the quality of POIs and POI responses similarly to how they consider other pieces of argumentative or responsive material in the rest of the debate. However, because failure to take a POI is indicative of a reduced level of engagement, judges must evaluate failure to take a POI as negatively impacting the persuasiveness of the speaker.

Cutting off a POI

Interrupting a debater who is giving a POI is known as 'cutting off'. POIs may be up to 15 seconds in length; however, a speaker may cut off a POI before 15 seconds and resume their own speech. Whenever a debater delivering a POI is cut off or their time elapses they must stop speaking, and sit down. If the person offering the POI does not stop speaking after 15 seconds, or after being cut off, the judge should intervene by calling 'order'.

If a POI is cut off before 15 seconds has elapsed, the judge should assess whether this cutting off was legitimate. If the POI was unreasonably cut off before the point could be clearly made, judges should treat it as if the speaker did not take a POI and apply the appropriate penalty. This is because speakers cannot meaningfully engage with POIs if they do not allow their opponents sufficient time in which to ask the POI.

Barracking/Badgering

After a POI has been offered to a speaker and rejected by them, another POI should not be offered within the next 15 seconds by any debater. Persistently breaching this rule (i.e., continuously offering POIs to a speaker in quick succession) is known as barracking or badgering. This is not permitted, as it is disruptive to the debate and unfair to the speaker.

POIs do not initiate a dialogue. Once the POI has been made/cut off, the debater making it sits down. They must wait the required time and offer a new POI if they wish to interrupt the current speaker again. The only exception to this is if the speaker was unable to catch the POI and ask the offeror to repeat or rephrase their question or comment. In this situation, the debater asking the POI may stay standing and repeat their question or comment.

Points of Clarification

Debaters sometimes offer POIs with the phrase 'Point of Clarification', usually to the PM's speech, to indicate that they wish to ask a question about how the PM is setting up the debate, rather than make an argument. This is permitted - but Points of Clarification otherwise function entirely as any other POI. Speakers are not obliged to take a POI just because it was labelled as a Point of Clarification. Taking a Point of Clarification does 'count' as taking a POI - because it is a POI. Points of Clarification have no special status in the rules whatsoever, speakers offering a POI are simply allowed a special exception to use the label 'Point of Clarification' when offering these types of POI.

Points of Clarification are POIs focused on clarifying the model; they should not be used to introduce substantive argument disguised as clarification.

1.4 Before the Debate

The Motion

Each round has a specific topic, known as the ‘motion’. The motions are set by a team of senior judges at the tournament known as the ‘Adjudication Core’ (also known as the ‘CA Team’, ‘CAP’ or ‘AdjCore’ for short). The CAP will announce the motion for each round of debates, along with the ‘draw’ (showing all the rooms in the tournament and the positions in which each team in the competition will be debating in each room) to all participants 15 minutes before the debates begin. If debaters are uncertain about the *literal meaning* of a word in the motion, they may ask a member of the CAP to define it for them. **They may not ask anyone other than a member of the CAP to explain any words in the motion, nor may they refer to online resources.** They may also not ask for any further assistance from the CAP beyond a simple definition of the word they are unfamiliar with.

Information, Context, or Definitions Accompanying Motions

On some occasions, the CAP may release an informational slide, known as an ‘infoslide’, ‘info-slide’, or ‘Information Slide’, prior to releasing the motion. This usually consists of a short explanatory paragraph which can serve several purposes, from simple clarifications of words in the motion to giving context and relevant information about potential issues in the debate.

Information provided in the infoslide should be assumed to be true for the purposes of the debate following it. For example, if the extra information comes in the form of a definition of a word or term in the motion, this definition should not be disputed in the round following it. However, teams are free to provide additional definitions, clarifications or contextual information during the debate, on top of whatever information is already provided within the infoslide.

Preparation Time

After the motion is released, teams have 15 minutes to prepare their speeches. **During these 15 minutes, the two speakers in a team must confer solely with each other while preparing.** Receiving assistance from anyone else during prep time, such as coaches, other members from their institutions, or judges, is strictly prohibited - teams spotted doing this should be reported, and may be penalised by disqualification from the tournament.¹ Teams must **not**, under any circumstances, use the Internet to research the motion, utilise generative artificial intelligence (AI) for any purpose, or to communicate with anyone that is not the CA team, the Organising Committee, or their partner. However, they may use their electronic devices as stopwatches, or as cameras to take photographs of the draw, motion and info-slide. They may also refer to electronic (offline) dictionaries. There are no exceptions unless teams receive authorisation in advance from the Equity team, as authorised by the CAP, due to special circumstances (such as

¹ We hope that no team at Worlds breaches these strict prohibitions. However, if you are a debater, and you witness another debater preparing with someone other than their partner or illegitimately using electronic devices, you should report this to a member of the CAP, or if they are not available, to any Chair judge or, if no Chair judges can be found, to any other judge. A judge informed about this should try to visually confirm that the team in question is indeed illegitimately preparing with outside assistance/illegitimately using electronic devices (ideally, they should also get another judge to witness this). They should then ask the team to provide their team name, and explain that preparing with someone other than your partner/using electronic devices for purposes other than timing or as an electronic dictionary is strictly prohibited. They should then (either immediately or after that round of debates is completed) inform a member of the CAP about the issue.

access needs). For the avoidance of doubt, teams that are allowed to use laptops are **not** allowed to use digital matter files, generative AI or online communications such as Google docs.

During the 15 minutes of preparation time, OG may prepare in the venue that will be used for their debate. Other teams, observers and judges should not enter the room until the preparation time is over.

Judges should call debaters into the debate room 15 minutes after the motion is announced. Teams must be ready to enter the debate room once the 15 minutes has elapsed. Late teams risk being replaced by a ‘swing team’ (a special *ad hoc* team created to replace them, which is not a fully participating team at the tournament), which will be summoned if they are not ready to enter the debate room after 15 minutes of preparation time. If the summoned swing team has reached the debate room, and the debate has begun, before the actual team has arrived, then the actual team will not be allowed to participate in the round, and will receive zero points for that round.

Pronoun Introductions

Before the debate begins, each of the participants in the room will be invited to introduce themselves and also be given the opportunity to introduce a gender pronoun.

There is no requirement to express a particular pronoun. Chairs should make this clear when they facilitate the introductions (of both speakers and adjudicators).

For example, chairs might say something like:

‘Before we start this debate, we will go around the room and introduce ourselves. At that time, you are welcome to state your pronoun preference if you wish to do so. If you do not want to state a pronoun, that is ok, and in that case everyone else please defer to gender neutral language’.

As a result, if you do not feel comfortable disclosing a pronoun or do not have a pronoun you wish to disclose, you may simply state your name (and speaker position) as your introduction.

If you do wish to state a gender pronoun, an example for doing so is:

‘Hello, my name is my gender pronoun is’

As the chair introduces each speaker, the chair can remind the room of the speaker’s pronoun (if applicable). For example, the chair might say:

‘I invite the Member of Government, X. Their pronouns are they/them.’

All participants should take note of the pronoun of each speaker and use that pronoun to refer to them (if applicable). **You should not assume anyone’s gender pronoun.**

If you mistakenly use the wrong pronoun, please apologise. Disregard for a person’s gender pronoun may be treated as an equity violation.

If a speaker or an adjudicator does not introduce a pronoun, **all other participants in the room should use gender neutral language** (e.g., ‘speaker’ or ‘Prime Minister’ or ‘adjudicator’).

1.5 Iron-personing

If, during any of the Preliminary Rounds, a member of a team is taken ill and requires medical treatment, or a recognised medical condition prohibits them from participating in a given Preliminary Round, the other member of the team is entitled to participate in the Preliminary Round as an 'iron-person' team. In an iron-person team, one speaker delivers both speeches. The speaker must prepare on their own. **In judging an iron-person team, the Adjudication Panel shall treat the team as if they were an ordinary team, and fill out the ballot accordingly (indicating that the team was an iron-person team on the ballot).**

The rules relating to iron-person teams shall operate at the discretion of the CAP and Equity Committee. Where there is a dispute between the two bodies regulating iron-person teams, the judgement of the Equity Committee shall take precedence.

From the perspective of other teams in the debate, and the judging panel, this team of one speaker giving both speeches functions just like a normal team - they may receive any rank in the debate from first to fourth, and will receive two speaker marks, one for each speech, and other teams in the debate will be awarded the other ranks as normal. In the 'tab' (the tabulated results for the tournament, maintained round on round and used to determine the break); however, the absent speaker will receive zero speaker points, and the iron-personing speaker will receive a single speaker score, the higher of the two speeches they gave. The iron-personing team may keep the team points that they received during the round, and these team points will be used to determine the draw for future rounds. Teams may still break as long as they are not speaking as an iron-person team for more than 3 preliminary rounds out of 9.

1.6 Breaches of Order

For the debate to be able to proceed properly, and for all speakers to have a fair chance to deliver their speeches, all debaters (and anyone else in the debate room) are required to refrain from disrupting the debate. Any of the following activities are considered to be disrupting the debate:

- Barracking/badgering
- Continuing to offer a POI after being cut off by the speaker speaking or by the Chair
- Offering POIs in any way other than those described in section 1.3 when not delivering a speech or a POI
- Speaking beyond 7 minutes with a 15 second grace period
- Talking in an audible volume or otherwise generating distracting noise during another speaker's speech
- Engaging in other highly distracting behaviour
- Using props (any physical object, diagram, etc.)
- Receiving any external communication during a given speech (e.g., notes passed to the speaker from their teammate)

These are not only breaches of the rules and/or appropriate debate conduct as it is commonly understood but are also breaches of *order*. Unlike other breaches of the rules (which simply damage a team's chance of getting a good result in the debate), breaches of order should be enforced by the Chair of the debate by *calling order*.

Calling Order

When the Chair of a debate utters 'order', it is a demand that all speakers immediately cease any of the breaches of order listed above. This should not happen often. Provided debaters adhere to the call to order, no further action is taken. A Chair should never call order for a breach of the rules which is not a breach of order.

Stopping the Clock

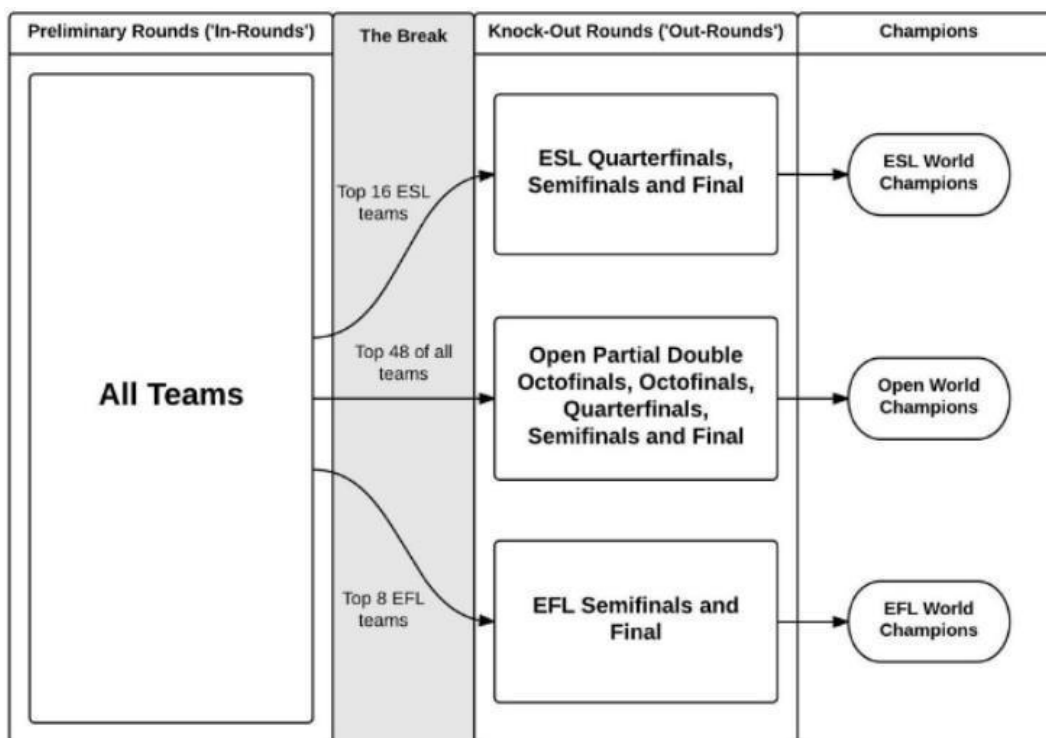
In **exceptional** circumstances, the Chair is entitled to clearly say 'stop the clock'; in which case the current speaker should immediately halt their speech, and the timekeeper of the debate should pause the stopwatch being used to time speeches. This measure should only be used in response to severe obstacles to the debate proceeding which need to be addressed urgently and cannot wait for the current speaker to finish their speech - for example, one of the debaters or judges fainting or suffering a medical emergency; or a severe and persistent disruption to the debate, such as a constantly heckling audience member, a technical failure in sound equipment that might be being used in the debate, and so forth.

In any such instance, the key objective of stopping the clock is to protect the welfare of all those involved in the debate, and to allow the obstacle to the debate proceeding to be dealt with as swiftly as possible (this may involve abandoning the use of any sound or recording equipment, having someone take an ill debater for medical attention, removing an unruly audience member from the room, and so forth). This will only very rarely be necessary in response to a breach of order, and is more commonly required due to an external interruption to the debate. Once this has been done, the Chair should check that the speaker is ready to resume the speech, call for the clock to be restarted, and allow the speaker to continue their speech from the point at which the clock was stopped.

1.7 Tournament Structure

The WUDC is structured in two halves (see the diagram below). The larger bulk of the tournament, usually taking place over the tournament's first three days, consists of a number of preliminary rounds (often termed 'in-rounds') in which all debaters at the tournament take part - historically there have been nine such in-rounds. Most of these rounds are 'open', meaning that teams find out the results of the debate, and receive feedback from judges, at the end of each round. The final few rounds, however, are 'closed' - results and feedback are not immediately given to speakers, but can be obtained from judges once the 'break' (see below) has been announced.²

² This postponement in giving the results ensures that teams do not arrive at the break with sure knowledge of whether they will advance to the knock-out stages or not.



After the in-rounds, the best performing teams in the tournament advance to a final set of knock-out rounds (often termed ‘out-rounds’) whilst the remaining teams do not - this process is known as ‘the break’. Teams are ranked in order according to the total ‘team points’ they have accumulated over the in-rounds (3 points for each first placed finish in a debate, 2 points for a second placed finish, 1 point for a third, and 0 points for a fourth), with teams tied on total team points ranked according to their total ‘speaker points’ (a mark out of 100 each speaker on the team receives for their speech in each room). At the current Worlds, 48 teams progress through to the ‘Open Break’ (for which any team at the tournament is eligible), usually 16 teams progress through to the ‘ESL Break’ (for which only teams with two ESL or EFL speakers are eligible) and usually 8 teams progress through to the ‘EFL Break’ (for which only teams with two EFL speakers are eligible). A team that is eligible to more than one break can indicate in advance to which break it wants to proceed (Open, ESL or EFL) - yet it cannot participate in more than one break. Those teams that make it into the three breaks then participate in three separate knock-out draws, progressing towards an Open Final, ESL Final and EFL Final, the winner of which becomes the World Champion in that category.³

³ This is assuming that the constitutional requirements for these breaks are met - the WUDC Constitution requires a minimum number of ESL and EFL eligible teams participate in the tournament for each stage of ESL or EFL finals to be included. If, for example, a smaller number of ESL teams are present, the ESL break may only be to Semifinals (eight teams); if a larger number are present, the break may be expanded to include Octofinals (32 teams).

2. Debating and Judging at WUDC

2.1 Winning a Debate

Teams in a debate are all aiming to win the debate. For both debaters and judges, the central statement on how teams win debates is as follows:

Teams win debates by being *persuasive* with respect to the *burdens* their side of the debate is attempting to prove, within the *constraints* set by the rules of BP Debating.

There are two important comments to make about this central statement:

- a) One could stand up in a debate and be persuasive about anything, but this will not help to win a debate unless it is relevant to the burdens teams are seeking to prove.
- b) The rules of debating constrain legitimate ways to be persuasive. For example, in the absence of rules, the OW could often be very persuasive by introducing entirely new arguments, but the rules prohibit this.

As such, elements of a speech can only help a team win a round if they are both persuasive and within the rules.

2.2 The ‘Ordinary Intelligent Voter’

In most walks of life, persuasiveness is highly subjective - the degree to which we are persuaded by something reflects our existing beliefs, our personal aesthetic or stylistic preferences, our particular interests, and so forth. It would be problematic if debating was judged so subjectively - outcomes would hinge as much on whom the judges were as on the debaters’ performance, with one side of the debate becoming much harder to win from because the judges were predisposed to disagree with it.

Consequently, as far as is humanly possible, **judges assess the persuasiveness of speeches according to a set of shared judging criteria, rather than according to their own views about the subject matter.** In particular, judges are asked to conceive of themselves as if they were a hypothetical ‘ordinary intelligent voter’ (sometimes also termed ‘average reasonable person’ or ‘informed global citizen’).

Facts, Knowledge, and Special Language

The ordinary intelligent voter has the sort of knowledge you’d expect from someone who regularly reads, but does not memorise, the front pages and world section of a major international newspaper (like the New York Times or the Economist) in the year leading up to WUDC. They do not read technical journals, specialist literature, or the like. They are, in short, a smart person who has a good deal of knowledge that is broad rather than deep. Imagine a bright

and reasonably well-read university student who is studying a subject completely alien to any topic that would help them understand the debate in question.

Debaters may certainly make reference to examples, facts and details the ordinary intelligent voter is not aware of, but they should explain rather than cite these examples, facts and details. While they may not know much on a specific topic by some debaters' standards, the ordinary intelligent voter is genuinely intelligent, and understands complex concepts, facts or arguments once they're explained. Where such examples are not explained beyond name-checking a country, judges should discount material they do understand that the ordinary intelligent voter would not. Judges should be bold in applying this rule: it is unfair on other teams in the room not to.

Importantly, the ordinary intelligent voter comes from nowhere, not where a particular judge comes from. So there are no 'domestic examples' requiring less explanation for the ordinary intelligent voter, even where everyone in the room comes from that country. Wherever you are from, assume your judges are from somewhere else.

Following on from the above, the ordinary intelligent voter does not know technical terms that one would require a particular university degree to understand. They can be assumed to possess the sort of generalist vocabulary that comes from a university education of some sort, but probably not from your specific degree. They do not have the sort of halfway-there economic or legal jargon that we as debaters have become familiar with either. Saying 'Laffer curve' to most people is equivalent to making some clever sounding noises. Similarly, using terms like 'economic efficiency' will lead to their being understood only as a layperson would grasp them, losing any technical specificity. Judges should judge accordingly and speakers who wish to make use of the extra specificity that technical terms convey should take the time to explain the connotations of the terms they wish to use.

Dispositions

This hypothetical ordinary intelligent voter doesn't have preformed views on the topic of the debate and isn't convinced by sophistry, deception or logical fallacies. They are open-minded and concerned to decide how to vote - they are thus willing to be convinced by the debaters who provide the most compelling case for or against a certain policy. The ordinary intelligent voter also does not have strong intuitions about the motion, and is thus open to being persuaded by a variety of intuitions. They do not judge debates based on their personal beliefs or political convictions, nor do they enter a debate thinking that one side is indefensible.

As described in the section above, they are well informed about political and social affairs but lack specialist knowledge. They are intelligent to the point of being able to understand and assess contrasting arguments (including sophisticated arguments) that are presented to them; but they keep themselves constrained to the material presented unless it patently contradicts common knowledge or is otherwise wildly implausible.

Judging as the Ordinary Intelligent Voter

As can perhaps already be intuited from the above paragraphs, the ordinary intelligent voter is quite unlike most, or perhaps any, real world people. But the concept of the 'ordinary intelligent voter' is a useful way of revealing a set of important characteristics that judges should aspire to in order to ensure that all teams receive a fair hearing in any debate. As such, the term 'ordinary

intelligent voter' describes the expectation that judges should:

- be aware of basic facts about the world (e.g., 'Syria is in the Middle East' would be considered basic);
- be familiar with issues and events that have made international headlines for a sustained period of time (e.g., judges should be aware of ageing population crises in developed countries. They should be expected to know that different countries had different models of response. For example, some countries increase immigration, while others increase the retirement age or implement pro-natalist policies. They do not necessarily need to be aware of the specifics of individual models each country has implemented unless a country's specific response has made international news headlines. For example, judges should be expected to know China abolished its one-child policy);
- avoid utilising personal knowledge that they have of the topic, unless it could reasonably be assumed to be held by someone who fulfils the previous two criteria;
- give little credit to appeals merely to emotion or authority, except where these have rational influence on an argument;
- avoid presuming a geographic, cultural, national, ethnic or other background when assessing arguments;
- avoid preferencing arguments or styles of speaking that match personal preferences;
- assess the merits of a proposed policy, solution or problem separate from any personal perspectives in relation to it.

This does not mean that speakers cannot make complex claims about complicated issues based on their own specialised knowledge, or indeed, that judges cannot be convinced by these claims. While judges should be assumed to have ordinary knowledge about various issues, they should also be fully capable of logically following and analysing a debate, and understanding complex concepts when explained. If teams wish to bring in their own specialised knowledge to the debate, they must be able to explain them in a way that is free of jargon and understandable by the ordinary intelligent voter.

Thinking as the ordinary intelligent voter does not absolve us from our responsibilities to actually judge the debate - to evaluate the logical flow of arguments, determine the extent to which teams have won them, and ensure that they have done so within the rules. We should not say 'while that was clearly irrational rabble-rousing, the ordinary intelligent voter would have fallen for it'. This not only leads to irrational conclusions, but also, generally, overestimates how much cleverer we are than an ordinary intelligent voter.

We emphasise that a key reason for judges to imagine themselves as the ordinary intelligent voter is to avoid relying on their subjective *tastes* as well as their subjective *beliefs*. Judges should remember that they are not aiming to evaluate who was cleverest, neatest or funniest, but who best used their cleverness, neatness and funniness to persuade us that the policy was a good or a bad idea.

2.3 Persuasiveness

Judges judge debates by assessing, without prejudice, which team in the debate was most

persuasive. The persuasiveness of an argument, in BP debating, is rooted in the plausible reasons that are offered to show that it is true and important (which we term ‘analysis’ or ‘matter’), and the clarity and rhetorical power with which these reasons are explained (which we term ‘style’ or ‘manner’).

It is crucial to understand that in BP debating, **analysis and style are not separate criteria on which an argument is assessed**. In particular, BP debating does not consider it possible for an argument to be persuasive *merely* because it was stylish. There is nothing persuasive in speaking a sentence clearly and powerfully if that sentence is not in fact a reason for an argument. And equally, reasons for an argument that cannot be understood by a judge cannot persuade them. Good style is about conveying a speaker’s analysis of arguments effectively to the judges. Style and analysis thus do not independently generate persuasiveness, but describe the necessary collective elements that make an argument persuasive. The fact that we discuss them, below, in separate sections should not detract from this.

Analysis

The analysis behind an argument consists of the reasons offered in support of it. Reasons can support arguments in a number of different ways, none of which is, in itself, ‘better’ or ‘more important’. Reasons might explain why arguments are true by:

- presenting empirical evidence for an argument;
- giving mechanistic links for why a certain outcome will come about;
- identifying widely shared moral intuitions in favour of an argument;
- exposing a damaging logical implication of a contrasting argument;
- identifying an emotive response that encourages us to care about a certain outcome;

...or doing various other things that encourage the ordinary intelligent voter to believe that an argument is true and important to the debate.

Reasons themselves may be stronger or weaker according to a number of important criteria, including:

- the precision of what the speaker says and;
- the detail with which relevant logical claims, empirical evidence, causal processes, moral intuitions, logical implications or other elements are explained.

Beyond these ways of identifying reasons within a speech that support arguments the speaker is making, judges deploy very minimal standards in assessing the degree of support a reason gives, whether the reason itself is plausible, and whether it therefore makes the speaker’s argument persuasive. Seriously implausible claims (such that any ordinary intelligent voter would not believe its logic and/or premises) provide weak, if any, support for an argument.

Certain things do not matter (in themselves) in evaluating how good a speaker’s analysis was:

- the number of arguments the speaker makes;
- how clever/innovative the argument was;
- how interesting the argument was;
- arguments that you’re aware of but which weren’t made.

What matters, once an argument is made, is how important its conclusion seems to be in the debate with respect to the burdens that each side is trying to prove, and the extent to which it seems to be analysed and responded to (and how well it withstood or was defended against such responses). Judges do not consider how important they thought a particular argument was, in the abstract, but rather how central it was to the overall contribution of any team or teams in this particular debate, and how strong the reasons speakers offered to support the claim that it was important/unimportant were.

Style

Arguments can be stylistically impressive in a range of ways - crucially, 'good style' should not be equated to 'the sort of style admired in my debating circuit/culture'. Speakers do not have 'bad style' because they don't speak with the particular idioms, mannerisms, coded references or established phrases used in the country their judge is from. Crucially, judges **must not** discredit arguments because of the style or accent in which they were delivered, and will be in breach of Equity should they do so.

Above all else, a 'strong accent' is not bad style. Everyone in the world has their own particular accent, and they all have their own accent strongly! When people talk about mild or strong accents, they mean how strong or mild the accent is compared to the accents with which they are familiar. This sort of subjective measuring is not a valid basis for judging certain styles as superior. There is only one legitimate way 'accent' can be a problem for a speaker at Worlds, and that is if judges genuinely cannot understand what the speaker is saying despite their *very best efforts to do so*. This is a problem in the same way that speaking too fast to be understood is a problem - judges have to understand the words a speaker says in order to evaluate them. This is a problem that could afflict *any* accent in principle - it is not just a problem for an 'ESL' or 'EFL' accent. Worlds is an international tournament, and speakers may find themselves judged by people from any nation. There is thus an obligation on all speakers to make themselves comprehensible to all judges and a burden on judges to do everything they can to understand a speaker's words and meaning. The tournament aims to be as inclusive as possible to speakers of all languages, but Worlds is inescapably an English-language-based competition. If judges cannot, despite their very best efforts, understand an argument, they cannot find it persuasive.

So, as suggested, one basic point underpins the judging of style at Worlds: there is wide global variation in what makes for an aesthetically pleasing style, **and subjective judgements of good style should not carry any weight in judging BP debating at an international tournament**. But this does not mean style is irrelevant. Worlds sets down a minimal number of principles to guide effective style that we take to be of fundamental and international applicability. As already noted, good style is about *conveying reasons effectively*.

Reasons are thus more compellingly delivered to the degree that:

- They are comprehensible.

As noted, the speaker's claims must be comprehensible to the judges to be evaluated. Technical jargon without explanation, speaking so fast you are incomprehensible, speaking so quietly you are not audible, slurring words, or fragmented sentences could all make an argument impossible to understand, and therefore could be unpersuasive. To be clear: judges

must make a dedicated effort to understand the speaker to the best of their ability, and **must not** automatically dismiss speeches as incomprehensible.

- They clearly and precisely convey the speaker's meaning.

Vagueness, ambiguity and confusing expressions necessarily make judges uncertain over the nature of the reasons the speaker is offering and how they support the speaker's argument. The more clearly and precisely speakers can convey their reasoning, the more persuasive it is.

- They effectively convey the emotional, moral, practical or other significance of the speaker's claim.

The key question a judge should ask themselves is: 'Is there additional information being conveyed via this stylistic choice?' If yes: then the rhetoric has amplified the effect of the logical analysis, and should be credited as making the argument more persuasive. If no: then the rhetoric has not been effective in conveying the significance of the logical analysis, and should not be credited as making the argument more persuasive.

Additional characterisation, illustration and framing that emphasise the logic being presented can all contribute to the persuasiveness of the argument. Word choice, phrasing, complexity of language, intonation, and other stylistic choices are not credited in isolation. They are only credited insofar as they **meaningfully** add to your arguments (for instance, using 'delta' to replace the word 'difference' does not meaningfully alter the content of your speech).

It is tempting but wrong to think that arguments in debating can be assessed through pure, cold, emotionless logic unaffected by language or tone. Making and assessing arguments is impossible unless one attaches a certain significance to outcomes, principles or claims, and appropriate use of language and tone can convey such significance.

It is crucial to note here that rhetoric cannot replace logical analysis - but rhetoric can amplify the effect of your logical analysis. Persuasive rhetoric does not necessarily need to be complex, so long as it communicates the significance of your point.

To reiterate: arguments cannot be persuasive just because they are stylish. Rather, style and analysis must work together to make an argument persuasive.

2.4 Contradictions

Teams (on either Government or Opposition) should not contradict themselves or their bench partners. Besides being unpersuasive, inconsistency is unfair to opposing teams. It cannot be reasonably expected from a debater to answer two contradicting lines of argumentation, especially if those are given in different times during the debate.

What Is a Contradiction?

A contradiction is: explicitly stating and taking a position opposite to one that is already made by your side; advancing claims that are mutually exclusive to the claims that have been advanced by your opening team, your partner, or earlier in your own speech.

A contradiction is not: a statement that is clearly pre-argumentative (i.e., a statement lacking the sufficient surrounding words to be a reason to support or not support the motion) or

mistakenly said. This is to avoid teams being unduly punished for a speaker mis-speaking and/or saying something otherwise inconsequential.

Contradictions within the Same Speech or within the Same Team

Teams cannot be credited for two mutually exclusive claims. They may only be credited for the first claim they have advanced. Subsequent claims which contradict or cannot coexist alongside the first claim should not be credited by the judges and opposing teams. This is due to the fact that internally inconsistent teams cannot simultaneously get credit for two areas of mutually exclusive argument.

In addition to not crediting contradictory claims, judges may also consider the extent to which the contradiction has undermined the strength of the team's arguments when determining the team's contribution to the debate. If either the speaker or the team directly contradicts themselves later in their speeches, this undermines their own points and should be taken into consideration during deliberation when determining how plausible their argument is, similar to as if an opposing team offered these arguments in refutation to the speaker. While the later claim should be disregarded, judges should evaluate how the contradiction affected the persuasiveness of the first claim.

However, judges should not credit opposing teams unless they point out the contradiction.

If a speaker clearly mis-speaks at the start of their speech, and they correct it afterwards, they should not have the rest of their speech discounted simply because it contradicts what they said first. Other than the instance of clear mis-speaking by the speaker, the argument made first should be considered to be the stance of the team, and later arguments that contradict the first argument should be discounted.

Contradictions between Teams on the Same Bench

It is important to note that contradictions or rebuttals of an Opening team's claims by their own Closing team should not be considered when determining the strength of Opening's arguments or their level of persuasiveness.

Arguments made by a Closing team that directly contradict their opening team's arguments should be ignored by the judge (i.e., the time spent by the Closing speaker contradicting their opening team, is equivalent to the speaker saying nothing at all).

This is to ensure that all teams in the debate are treated fairly, as Closing teams have a rules-based obligation to stay consistent with their opening teams. This also ensures that debates are coherent and that teams are not forced to defend opposing claims or respond to contradictory cases.

Making an 'even if' argument (along the lines of 'even if OO were wrong about this, we're going to show that this motion should still be defeated') does not constitute knifing. However, as with any other extension, an 'even if' extension will not provide good grounds for a Closing team win unless it improves the bench's persuasive position.

How Teams Should Deal with Contradictions from the Other Side

It is good practice for teams to point out contradictions (if they exist) in the other side's case, including between the two teams on the opposing bench. Whenever there is a contradiction, teams should consider the first claim to be the version they must engage with.

By way of example: OG offers 2 claims which are mutually exclusive - claims A & B. Claim B (the later one) should not be credited by judges when evaluating the contribution of OG, however claim B can still weaken the persuasiveness of claim A. When engaging with OG, other teams should consider claim A to be the line of argumentation OG pursues - i.e., Opposition should respond to it and CG must be consistent with it (see section 2.9 for issues relating to definitional challenges).

2.5 Rebuttal, Engagement, and Comparisons

The outcome of the debate should depend on what the teams say. Judges must not intervene in the debate. Do not invent arguments for teams, do not complete arguments, and do not rebut arguments. Engagement from closing half teams should not benefit their opening (e.g., closing half rebuttal should not influence the pairwise comparison between OG and OO).

We do not automatically dismiss arguments just because we disagree, or because we can see weaknesses in them. Arguments are persuasive and impactful once they are made and substantiated; they become less persuasive and impactful if they are contradicted internally, or responded to by other teams.

This has an important implication: if OG, for instance, makes arguments where the conclusion is 'we should do the policy,' and every other team ignores those arguments, then OG does not lose because 'the debate moved on from them'. Rather, their unrebutted arguments should still be treated as impactful and should be weighed as such. That does NOT mean that the unresponded-to arguments have a particular effect on the ranking of OG in this example. Judges still need to consider how significant an argument is before deciding how it affects the ranking of teams in a debate.

If an argument is clearly absurd (such that you cannot conceive of any ordinary intelligent voter believing its logic and/or premises), or it was of marginal importance to the speech of the speaker making the argument, then it is reasonable for a responding team to decide to spend their time elsewhere, particularly where there is other stronger material in the round. Furthermore, judges are entitled to assess how well substantiated an argument is - an argument that is just an assertion ('as we'd all agree, language constructs reality') without any subsequent substantiation should not receive much credit. There is no absolute duty for a speaker to 'hit every argument' from the other side. However, it may be advantageous for other teams to point out and respond to weakly constructed arguments. If major claims go unchallenged by teams, this should be counted as conceded by the team which has passed up the opportunity to respond.

Rebuttal consists of any material offered by a speaker which demonstrates why arguments offered by teams on the other side of the debate are wrong, irrelevant, comparatively unimportant, insufficient, inadequate, or otherwise inferior to the contributions of the speaker's own side of the debate. Rebuttal need not be explicitly labelled 'rebuttal' (though it may be sensible for speakers to do so), and it may occur at the beginning, end, middle or through the entirety of a speech. Material labelled as rebuttal can be constructive as well as rebuttal, and material labelled as constructive can also function as rebuttal. **Rebuttal does not, therefore, denote**

some special sort of argument or analysis - it simply refers to any material that engages directly with arguments raised by the other side.

Where teams have a chance to rebut each other, assessing relative contributions in this way is easy. Judges should track the argument and assess, given their responses to each other, which team's contribution was more significant in furthering their cause to logically persuade us that we should do the policy, or that we should not.

But where teams don't get a chance to rebut others, determining who was more persuasive is trickier. This happens fairly often, for example:

- between teams on diagonals
- when the OW explains something in a new way
- when opening teams are shut out of POIs

In these circumstances judges are forced to perform some more independent assessment of the arguments made. Judges will have to assess not only which arguments are most important, but also which ones are most clearly proven. Arguments that require the judge to make numerous logical leaps are better than no arguments at all but are not preferable to a well-reasoned argument that rests on fewer unsubstantiated assumptions.

Assessing arguments will also involve a comparison with existing material within the debate. For instance, when judges compare two teams on a diagonal (e.g., OG and CO), they should first ask whether anything in the earlier-speaking team's case is inherently responsive. Did the opening team preempt any material within their case construction or their substantives? Did the later-speaking team being assessed deal with the stronger parts of the opening team's case, or merely the weaker parts? Check whether they allowed the diagonal team in on POIs, to give them an opportunity to engage. Deliberately shutting out engagement from a team whose material is relevant is often obvious and very unpersuasive.

Crucially, in these instances, judges must not make new arguments for top-half teams; they may only reasonably interpret and apply existing contributions. Closing teams are not required to explicitly justify why their arguments should be valued more than those from the Opening teams in order to win over Opening teams. While 'bench weighing' can increase the persuasive value of a team's arguments, the absence of 'bench weighing' is not, by itself, a sufficient reason for a team to lose. If a team provides 'bench weighing', the judging panel should assess it with the same rigour as any other argument in the debate, rather than accepting it without question. If the Closing team does not provide an explicit justification, they should not be penalised. Judges are still expected to independently weigh the arguments when deciding which opening or Closing team wins. For further detail on weighing, see 2.5 on weighing competing frameworks.

2.6 Burdens

As stated earlier, there is no value in being persuasive about an argument that is irrelevant to the debate. In assessing what contributions are relevant, it is helpful to consider the 'burdens' a team has to meet in the debate.

Burdens on teams cannot be created simply by another team asserting that they exist, and judges

should not accept these assertions if they are not backed up by analysis. Teams and judges should not push unrequired burdens onto their opponents. Moreover, even if a team fails to meet a burden, that does not mean that they automatically lose the debate. Judges should consider analysis advanced by teams even if that analysis does not necessarily meet the burdens in the debate.

However, there are two key ways that a burden can legitimately be attributed to a team (and speakers may legitimately point out such burdens, and explain why they or other teams need to meet them).

First, a burden may be implied by the motion itself. If, for example, the motion is ‘**This House Would prioritise the vaccination of law-abiding citizens in the case of major epidemics**’, government teams need to demonstrate that *in major epidemics* the vaccination of *law-abiding citizens* should be *prioritised*. Government teams do not need to demonstrate that vaccinations of law-abiding citizens should be prioritised in general (outside of major epidemics), or that *only* law-abiding citizens should be vaccinated (law-abiding citizens should simply be prioritised). The way OG defines the motion (see below) may affect these burdens, however. Opposition teams need to demonstrate that the Government is wrong: that the policy of prioritising law-abiding citizens for vaccination in major epidemics should be opposed. They do not necessarily need to show that law-abiding citizens should not be prioritised in any way under any conditions (though the fact that we do prioritise law-abiding citizens in other cases might be used as evidence of a principle that supports prioritising law-abiding citizens in this case).

Second, burdens can also be set by specific analysis teams take up. For example, if the motion is ‘**This House Believes That assassination is a legitimate tool of foreign policy**’, the LO may initially argue that assassination is a severe breach of international law. For this to be relevant to the debate, OO has a burden to show that illegality matters for illegitimacy. This burden is especially strong if the DPM then states that they accept that assassination is illegal, but argues that illegality is a poor basis for believing an act illegitimate. Unless Opposition teams now provide superior reasons to think that the illegality of an act under international law *is* a reason to deem it illegitimate, it is not relevant to the burdens they need to prove to merely keep pointing out that assassination is illegal, or provide more detail on how it is illegal. Both sides now agree that assassination is illegal, and continuing to agree with this achieves nothing. What the sides now disagree on is the implications this has for assassination’s legitimacy, and it is this which they have a burden to prove.

Weighing Competing Frameworks

As evidenced by the above examples, teams will often dispute the criteria by which the round should be adjudicated on, and argue that points should be judged according to certain frameworks and standards. This is permitted: teams are allowed to debate what criteria should be used to assess whether a policy is good as part of arguing that it is, in fact, good.

Judges should adjudicate this debate about criteria - **they should not just apply their own preferred criteria**. They should adjudicate this on the following basis:

- Is there one criterion or principle that all teams explicitly agree is true and important (e.g., all teams explicitly agree that their goal is to save the most number of lives, and the debate is about the best way to do so)?

- If not, is there one criterion or principle that all teams implicitly agree is true and important (e.g., while no team explicitly articulates that their goal is to save the most number of lives, all the analysis advanced by teams points in this direction)?
- If not, is there one criterion or principle that one team in the round has successfully proven to be true and important (i.e., If no team agrees on one criterion, and all teams are asserting different metrics, which team has provided the best reasons to believe that their metric is the most important one in the round)?
- Where none of these apply, judge based on what the Ordinary Intelligent Voter would take to be important. This should be a last resort measure only, as it is very rare that none of the aforementioned scenarios would apply.

One common form of this mistake is to assume a utilitarian ('what leads to the best consequences') framework. This should not be assumed without a team presenting supporting arguments for doing so. It is also wrong to disregard principled argumentation explaining that particular effects are more important than others for reasons unconnected with utility maximisation. So, judges should listen to teams' arguments about what our aims and principles should be, and evaluate the claims of harms or benefits in that context. This can make the claims about *how* we should determine the right policy particularly vital, and they may fundamentally reshape team's burdens in the debate.

For example, if in the debate '**This House Would invade North Korea**' Opposition teams successfully prove that 'war is always wrong, regardless of the practical benefits' (they must do more than assert it), Government teams will likely now need to offer reasons to believe that a practical calculus is relevant if they want to advance purely practical reasons in favour of the invasion.

Judges should generally be wary of considering an argument **completely** irrelevant because of a principled framework advocated by their opponents. It is very unlikely that any team will ever prove their view of the appropriate criteria to be completely and undeniably true and that, consequently, arguments which do not fit those criteria should be completely dismissed out of hand. **It is thus often more appropriate to treat arguments as less persuasive when they rest on criteria which another team has suggested are not relevant, rather than ruling them out completely.**

2.7 Motion Types

Motions can come in a few different guises, often hinted at by the words used to introduce the motion ('**This House Would...**', '**This House Believes That...**', '**This House Supports...**') and again, this can affect the burdens teams face. The rules for debating each motion type are detailed below.

Policy Motions

This House Would ('THW')

Motions of the form '**This House Would [do X]**' involve Government teams arguing that they should be enacting policy X. A policy is a concrete course of action that Government teams wish to convince the judges should be implemented. Such motions are about whether the House should

do X - with Government teams arguing that they should and Opposition teams arguing that they should not.

These debates are purely normative. They do not require teams to discuss whether or not policy X is likely to be enacted in the real world, or whether or not policy X is currently status quo.

For the purposes of the debate, the Government teams have the powers of the actor implied by the perspective of the motion (this may be a government, an individual, society, etc.), and the debate is about whether they should or should not do a policy or action, not whether their real world counterparts will or will not. It should be assumed that the policy will be implemented in the manner that the Government teams set up (also known as Government fiat). As such, it is never a valid line of opposition to such motions to state that ‘but the government would never do this’ or, more subtly, ‘but politicians would never pass this law’.

Take, for example, the motion ‘**THW ban cigarettes**’. The debate should assume that the Government team has the power to implement such a policy and that this policy will therefore pass the approval of Congress or Parliament; however, the Government team cannot control reactions to this policy, and cannot assume that everyone will behave in a compliant manner once the policy is passed. The question of the debate is whether or not the policy should be enacted in the manner that the Government team has set out, not just about whether or not cigarettes are good or bad. It is perfectly possible for the Opposition team to agree that cigarettes are bad, but oppose the policy of banning cigarettes altogether.

For Policy motions, Opposition teams may choose to defend status quo, propose a counter-proposition or suggest an alternative (see section 2.10). It is not necessary for Opposition teams to present a counter-proposition, though it may be beneficial in some instances.

If presenting a counter-proposition, Opposition teams are granted the same amount of fiat power that Government teams have: the debate should assume that whatever counter-proposition Opposition proposes will also be implemented, and it would be similarly futile to argue that Opposition’s counter-proposition would never be passed by any parliament in real life. However, it is crucial to note that the Opposition’s counter-proposition must use the same or fewer resources than the Government’s policy. More information can be found in the section on counter-propping.

This House Believes That X should (‘**THBT X should**’)

Motions of the form ‘**THBT [X] should [Y]**’ where [X] is an actor and [Y] is an action or policy are policy motions. Even though these motions are phrased as statements of belief, government teams have fiat to model how X is done, and opposition teams have fiat to propose a counter-proposition (see section 2.10 on counterpropping). However, this fiat is from the perspective of the actor [X], which means government and opposition teams may only fiat in things which the actor has the power to do. These motions are about whether or not the statement is true from the perspective of a neutral observer. They should not be confused with **actor motions** or **analysis motions**, discussed in the sections below.

Analysis Motions

This House Believes That (‘**THBT**’)

Motions that begin with ‘**This House Believes That [X]**’ are value judgement debates about the

statement represented by X. They require Government teams to argue that X is true, whilst Opposition teams argue that X is not true. Opening Government may not implement a model in these debates, as they are not proposing a policy.

Take, for example, the motion **‘THBT there is no moral obligation to follow the law’**. The debate is about whether or not the statement is true, *not* about whether or not the government should do anything about the statement (by, for instance, abolishing all laws, which is in any case implausible). Government teams need not have a model; they should, however, still define terms within the debate. In this case, they should define what a moral obligation is.

This House Supports/Opposes (‘THS/O’)

Motions that begin with **‘This House Supports/Opposes [X]’** are debates that require teams to support or oppose [X] in the way it is likely to manifest in the world. The burden on Government is to prove that [X] does or will do more good than harm in ‘This House supports’ motions (or more harm than good, in a ‘This House opposes’ motion). Similarly, the burden on Opposition is to prove that [X] does or will do more harm than good in ‘This House supports’ motions (or more good than harm in a ‘This House opposes’ motion). Both the Government and Opposition teams do not have the fiat to assert or dictate where, when, or how [X] will manifest. Instead, the debate should be evaluated based on the probable outcomes of [X], as analysed or agreed upon by the teams participating in the debate.

Take, for example, the motion **‘THS US involvement in the Middle East’**. Government teams must argue that US involvement is positive in totality, without picking and choosing which aspects of this motion they are supporting. Similarly, Opposition teams must oppose this motion in totality, without picking and choosing what to oppose. Teams cannot support only favourable aspects of US involvement, nor can they oppose only unfavourable aspects of US involvement.

Secondly, Government cannot model how US involvement will occur. They can argue that US involvement is likely to happen in a certain way, but this characterization is open to challenge by the other teams. To reiterate, neither Government nor Opposition have fiat power in This House Supports / Opposes debates.

This House Prefers (‘THP’)

Motions that begin with **‘This House Prefers’** function in the same way as other analytical debates, with one important difference: Opposition teams are bound to defend the specific comparison provided by the motion. They must either:

- In motions phrased **THP X to Y**: defend Y or
- In motions phrased **THP X**: defend status quo.

For example, in the motion **‘THP the conscription by lottery as a means of enrolling people in the army to aggressive recruitment of volunteers’**, Government must argue that conscription by lottery is preferable to aggressive recruitment of volunteers, and Opposition must argue that aggressive recruitment of volunteers is better than conscription by lottery.

In the motion **‘THP conscription by lottery as a means of enrolling people in the army’**, Government teams must argue in favour of conscription by lottery. Opposition must argue in

favour of army enrolment policies **as they are in the status quo**. They **cannot** argue in favour of abolishing the army, or reducing enrolment in the army (remember, Opposition is bound to defend the specific comparison in the motion!).

Debaters should be aware that there is a unique version of THP motions, which are phrased **‘THP a world in which X’**. These types of motion set a burden on Government teams to envision and argue in favour of the alternate world described in the motion. As in all other types of THP motions, Opposition is still bound to defend the status quo, or whatever comparison is presented in the motion.

In the motion **‘THP a world in which organised religion does not exist’**, Government teams need to conceptualise an alternative world without organised religion. This motion is also backwards looking: it requires teams to consider how the world would have developed had organised religion never existed. Here, it is reasonable to expect the debate to contain some discussion of how the trajectory of human history or development would have been impacted. This is similar to how counterfactuals work in THR motions (see section below on Regrets motions).

As these debates require the conceptualisation of an alternative world, arguments about transitions between the status quo and the alternative world are not permissible. For example, using the previous motion, teams should not discuss backlash from the demise of organised religion, as organised religion would have never existed in this alternative world.

In ‘THP a world’ debates, the counterfactual world exists by fiat. Teams do not have to explain how the counterfactual world came about. ‘Retrocausality arguments’, or arguments about the process by which the counterfactual world came into being, are not permissible and should not be credited. For example, using the previous motion, Opposition teams cannot argue that this world could only have resulted in governments actively suppressing attempts to start religion by killing everyone involved.

Debaters should also use their common sense to determine the point at which this new world most likely diverged from the status quo. For example, some motions mention the introduction of a new technology. It would, in most cases, be unreasonable for teams to assume that this technology existed 2000 years ago. It would be more reasonable to assume that this technology was recently introduced. Similarly, in the motion **‘THP a world where Kamala Harris won the US Presidential Election’**, it should be clear that teams are meant to discuss the election of 2024, when she was the Democratic nominee, and not, for instance, the election of 1800, or even the election of 2020.

Motions which use the words to the effect of ‘prefer’ (e.g., ‘rather than’, ‘as opposed to’, ‘instead of’) would also **bind Opposition teams to defending the specific comparison**.

For example, if the motion is **‘THBT it is in the interest of Vietnam to strengthen relations with the US rather than China’**, Opposition teams are bound to defending strengthening relations with China. Opposition teams would not be allowed to defend ‘we do both and play the countries off against each other’. If the motion is **‘THBT it is in the interest of Vietnam to strengthen relations with the US’**, Opposition teams have to defend the status quo. If the status quo involves playing the countries off against each other, then the Opposition may defend that comparative.

This House Regrets (‘THR’)

Motions that begin with **‘This House Regrets [X]’** ask whether the absence of X would have been good for the world. With the benefit of hindsight, teams can consider past harms and benefits. Teams can also consider how a world without the existence of X may prevent future harms or create future benefits. Teams should describe how an alternative world that developed without X occurring would look like. This is also known as a ‘counterfactual’. For example, with the motion **‘This House Regrets the selection of Kamala Harris as the Democratic nominee’**, teams cannot just debate the merits or downsides of Harris as a Democratic nominee. Instead, they should consider who alternative Democratic nominees might have been, and whether those nominees would have led to better or worse political outcomes than in the status quo.

This House Predicts

Motions that begin with **‘This House Predicts that [X]’** ask teams to analytically prove that X will happen. This motion requires teams to prove that X will happen, in the same way that an analysis motion would ask them to prove that X is true. There is no burden on teams to prove that the outcome of X is good or morally desirable. For example, with the motion **‘This House Predicts that Germany will not meet its climate goals’**, teams should not debate whether it would be good if Germany was to meet its climate goals or if it has a duty to do so, but rather whether or not, given what we know about Germany at the time the motion is set, we believe they will hit their goal. This means judges need to evaluate the level of persuasiveness teams put forward as to whether it is likely that [X] will happen, rather than whether the outcome of [X] is desirable.

Motions that are phrased as **‘THBT [X] will happen’**, **‘THBT [X] will occur’** or similar are analytically equivalent to **‘This House Predicts that [X]’**.

CAPs should not use the abbreviation ‘THP’ for this motion type because ‘THP’ is presumed to refer to **‘This House Prefers ...’**.

Actor Motions

Motions that begin with **‘This House, as [A], would do [X]’** are **actor motions**. Actor motions require teams to consider the motion from the actor’s own perspective instead of merely considering what would be best for the world. This means focusing on the actor’s own reasoned judgement about what they ought to do.

What an actor should do is different to what the actor is likely to do. Actor motions are not about predicting likely behaviour but rather about what most conforms with the values, interests, and duties of the actor in question. Whilst past statements of intent help us understand an actor’s perspective, the actor may nevertheless be persuaded to follow a different path.

Persuading an actor & conflicting reasons

In reasoning about a course of action, an actor will be swayed by various kinds of reasons. Any actor will give due weight to their values, interests, and duties in deciding how to proceed. It is common for these reasons to conflict with each other, such as when one’s duties require them to act against their self-interest. To resolve such conflicts between reasons, judges should consider the actor’s own value system to determine which reason will ultimately be stronger in persuading the actor.

Importantly, it should not be assumed that actors are purely self-interested, nor is it true that actors solely care about maximising pleasure and minimising pain. For these arguments to be persuasive, teams must give reasons that the actor should pursue selfish or pleasure-maximising behaviour, even at the expense of their other values.

What does it mean to ‘adopt an actor’s perspective’?

Adopting an actor’s perspective requires us to defer to the actor’s values, and priorities between those values, even if the ordinary intelligent voter would not agree. For the purposes of the debate, judges should hold those values with the same intensity as the actor and consider whether the actor would be swayed to change their beliefs as a result of the arguments. Actors hold some values more firmly than others, and are unlikely to be swayed to act against the values they most strongly believe in.

Crucially, adopting the actor’s perspective does not include any limitations in the actor’s knowledge of the world. Any information that teams give is information the actor becomes aware of, though whether or not the actor accepts or rejects this information is within the debate. Whether the actor accepts or rejects information contradictory to their existing beliefs depends on the strength of teams’ analysis supporting its truth, and the strength of those beliefs. Teams may not make arguments premised on the actor being ignorant of a fact.

When an actor’s existing beliefs are challenged, judges should consider whether the actor, upon reasoned reflection, would continue to hold those beliefs after considering the arguments made by teams. The actor should be assumed to be fully capable of reasoning, and to be operating free from manipulation. Judges should not credit arguments that would only be persuasive to an actor whose reasoning is erroneous (e.g., presenting the actor with a false dilemma).

Alternative formulations of actor motions

A motion worded ‘This House Would’ should be treated as an actor motion if it contains an Information Slide describing the perspective of an actor (commonly starting with the words ‘You are a ...’).

Some actor motions will be phrased as ‘This House, as [A], supports [X]/believes that [X]/prefers [X]’ etc. These motions are subject to the same rules as the corresponding analysis motions for the purposes of setting up the debate. However, arguments are evaluated under the rules applying to actor motions.

Groups of people as actors

Where an actor motion refers to a group of people (for instance a country, a business or an organisation) it may be unclear what set of people are included within that group. In such a case, the OG team has definitional fiat over who is included within that group, and so long as that definition is reasonable, it should be accepted by judges and other teams. For example, on the motion ‘**This House, as Iran, supports a nuclear deal with the United States**’ the OG team has the power to define whether ‘Iran’ in this instance refers to the government of Iran, the people of Iran as a whole, or another reasonable definition. So long as this definition does not unreasonably narrow the debate (see section 2.9 under ‘squirreling’), this definition should be accepted in the round.

We recommend that CA teams avoid setting actor motions with ambiguous actors, and instead opt to explicitly clarify actors to the extent possible (e.g., setting **‘This House, as the Iranian government, supports a nuclear deal with the United States’** or **‘This House, as the Iranian public, supports a nuclear deal with the United States’** instead) in order to avoid messy debates due to definitional ambiguity.

This does not mean that impacts on other stakeholders cannot be argued about, however the amount that the judges ought to care about impacts on other stakeholders is dependent on how much teams can show that those impacts on other stakeholders are cared about by the people included within the actor in the motion.

Alternative Motions

Some motion formulas are not defined in the manual. **Any currently undefined phrasings of motions should be treated as ‘This House Supports [X]’ or ‘This House Opposes [X]’, where [X] is the core element of the motion.** Whether it becomes supports or opposes should be decided by which is semantically closer to the meaning of the original motion wording.

This most commonly involves non-standard verbs being used in place of ‘supports’ or ‘opposes’. For instance, in the motion ‘This house celebrates X’, ‘celebrates’ is clearly a positive word closer to ‘supports’ than ‘opposes’, so this should be treated as ‘This house supports X’.

Where there is a lack of clarity over how the motion can be transformed into a ‘supports’ or ‘opposes’ debate, OG has the definitional fiat to define the nature of the motion, so long as their definition is reasonable. This can include a lack of clarity over either what the core element (the [X] which is supported or opposed) is, or over whether it is more reasonably a ‘supports’ or ‘opposes’ debate. For instance, the motion **‘TH sanctions Milorad Dodik’** may be reasonably read multiple ways (sanction either meaning ‘to allow’ - closer to support - or referring to a set of punishments imposed for bad behaviour - closer to oppose), which allows OG to use their definitional fiat to define the debate. We encourage CA teams to avoid setting motions with such ambiguous formulations in order to avoid messy debates.

2.8 Role Fulfilment

Role fulfilment, in brief, is the name given to the particular duties given to each team on the table, arising because of their particular position, beyond the general duty to ‘make persuasive arguments’. Some such duties exist to ensure fairness by specifying additional constraints on the debaters to reflect the idiosyncrasies of BP debating as a method of persuading an ordinary intelligent voter.

A debater who gives an excellent fifteen-minute speech, or submits a persuasive essay or a set of visual aids, will not be entitled to credit for doing so, regardless of how persuasive these would have been in conveying their reasons for affirming or rejecting the motion. Doing so involves breaking the rules, and cannot entitle them to credit. Role fulfilment is a necessary (but not sufficient!) condition for a team to make persuasive arguments. Incomplete role fulfilment should not be used as the sole justification to fourth teams.

The duties associated with role fulfilment are as follows:

- For the PM, to ensure the debate is adequately defined (see section 2.9).
- For the Member speakers (both Government and Opposition), to extend the debate (explained in section 2.8).
- For all speakers, to ensure that their arguments are consistent with all other arguments made by themselves, their teammates, and the other team on their side of the debate (see section 2.4 for more information on contradictions).
- For all speakers, to take **at least one POI** during their speeches and to offer POIs on a regular basis (see section 1.4 on POIs).
- For all speakers, to speak within the time frame allotted (see section 1.2 on speech timing).

We emphasise here that there is no such thing as an ‘automatic fourth’ or any automatic penalty for a failure to comply with the rules in this document. A team that breaches an element of role fulfilment may still be sufficiently persuasive to beat other teams in the debate; particularly, but not exclusively, when multiple teams in the debate have role fulfilment issues.

2.9 Definitions and Models

A model refers to OG’s explanation of **how** the policy they are proposing will be implemented. If the motion requires a model, or if OG wishes to propose a model, this must be explained in the PM speech. The DPM may clarify parts of the model in response to any confusion by the Opposition teams, but should not introduce a new model or new substantive portions of the model. Government teams are allowed a level of fiat in proposing their model - this is explained in more detail in section 2.3, under ‘Policy Motions’.

Regardless of motion type, OG is tasked with defining the debate. This is known as ‘**definitional fiat**’ and gives OG the right to define what the words in the motion mean. This is distinct from ‘**modelling fiat**’, which OG has in certain motion types (namely policy debates) to define how the policy will be implemented. Debates are about the motion as defined by OG, not about what other debaters or judges in the room thought the words in the motion meant. This often includes defining ambiguous words or concepts in the motion for the purposes of the debate. The definition forms the subject matter of the debate. These definitions are not subject to ‘likelihood analysis’ (OG does not need to prove the likelihood of their definition, and other teams cannot disprove the definition by showing it is unlikely). Definitions can only be altered by formal definition challenges, as outlined below. If the motion ‘**This House Would privatise education**’ is defined as ‘making all universities independent companies’ (a fair definition), then that is what the debate is about for the remainder of the eight speeches.

The PM should ensure the debate is adequately defined. For example, some debates will use words like ‘widespread use of X’. The term ‘widespread use of X’ refers to a state of the world where X is a relatively broad phenomenon. The PM has the ability to define what degree of use qualifies as ‘widespread’ as long as it is reasonable. However, the PM does not have power to define the manner of use of X, which must be characterised through likelihood analysis.

Some motions do not require any explicit definitional work from the Prime Minister, as the words in the motion are self-explanatory for the purposes of the debate.

The definition should be at the level of generality implied by the motion. It is legitimate for OG to exclude marginal and extreme cases (‘we’re banning cosmetic surgery like the motion says, but not for burns victims’). It is not legitimate to include only marginal and extreme examples

(‘we’re banning cosmetic surgery like the motion says, but only for children’). If CAPs wish a debate to be narrowed down in some specific and radical way, they will state this in the motion. To give another example, if the motion is **‘This House Would use community service as a punishment in place of prisons’**, and the Government bench states that it will only do this for young non-violent offenders, this is a severe and invalid restriction of the motion, excluding the considerable majority of cases to which a literal reading of the motion (which mentioned no limits to specific categories of prisoner) would seem to apply.

If teams wish to exclude non-marginal cases from the debate, they must provide a clear criteria on which cases are excluded and a compelling justification for doing so, and their exclusions should not unfairly disadvantage other teams in the debate. Common forms of legitimate restriction include explicitly limiting or focusing the debate onto broad sets of cases where the motion seems particularly applicable or would most plausibly be implemented.

For example, Government teams might argue that the scope of the debate is most relevant to countries in the developing world, and provide reasons for suggesting this. This is not to say that impacts on countries in the developed world are considered out of the debate - merely that a team has provided reasons why the debate might plausibly focus on a particular area. Again, the question in all cases is one of fairness and consistency with the original motion. This is ascertained by asking whether the definition excludes a large number of cases to which the motion seems to apply, and in doing so unbalances the debate. If not, the definition is likely to be legitimate. Still, as a general rule, it is sensible for OG teams to avoid restricting and limiting motions too much.

The definition should not be restricted to a specific time or place. Unless the motion specifies otherwise, it should be assumed to apply to the bulk of the world’s states. Some motions may presume a certain level of state capacity: for instance, the motion **‘This House Would allow citizens to sell their votes to others’** will only be relevant in states that are minimally democratic, and Opening Government may specify this without being accused of place-setting. However, teams may not restrict the motion to a particular location (for instance, they may not limit the debate to just the United States of America, or European capital cities).

Unless the motion specifies a particular time, OG should define the debate as being set in the present day. It is invalid for OG to define the debate as being in some particular time. For example, if the motion is **‘This House Would allow abortion’**, OG cannot define the debate as being about whether the judges in the key US case of *Roe v Wade* should have reached the decision they did at the time of that case. However, proposing a specific time scale for a motion does not constitute time-setting provided it keeps implementation reasonably close to the present day. So saying ‘we will allow a two year transition period for businesses to adapt to the proposed changes our policy creates before we proceed to full implementation’ is legitimate, whereas saying ‘we believe this policy should eventually be implemented, perhaps in one or two decades, once all countries will have fully harmonised to its requirements’ is not. Another example is debates which use the term ‘rise of’. **The term ‘rise of’ refers to the development of a trend up until the debate. It does not refer to the future growth of the trend.**

Squirrelling

A definition may also be invalid if it is a ‘squirrel’. A ‘squirrel’ is a definition of the motion which seeks to diminish or evade the burden of proof the motion places on OG.

A definition may be considered a ‘squirrel’ if it is literally inconsistent with the words of the actual motion that was set. If, for example, the motion is ‘**This House Would place tolls on all roads**’ and OG suggests they would place tolls only on major motorways, this is clearly invalid, since the motion specifically says ‘all roads’.

A definition may also be considered a ‘squirrel’ if it is not debatable. For instance, if the motion is ‘**THR the use of feminist messaging in commercials**’, it would be illegitimate to claim that this debate is only about negative instances of feminist messaging in commercials as this unfairly limits the scope of the debate.

If teams make arguments purely based on a squirrel, and their squirrel is challenged, then their arguments may be called into question as well. However, if teams make arguments that may apply to both the squirrelled definition as well as a legitimate definition, then their arguments should be judged based on the content of the argument.

Vague definitions

A vague definition does not clearly answer certain vital questions about what is meant by the motion or what will happen under the policy OG is defending. A definition cannot specify everything and OG is not expected to be exhaustive. But common points of vagueness include, where the debate requires it to function fairly, failing to specify: exactly what groups of people a policy applies to, the circumstances where it will be implemented, the agent who will implement the policy, or the consequences for those who resist or defy it.

A definition can be vague to different degrees. Crucially, a vague definition is not an invalid definition - it just undermines the persuasiveness of OG to the degree that it is unclear exactly what they are proposing to do. The proper response from Opposition teams is to identify this vagueness and its impact on the debate, via POIs or in their speeches. Later government speakers can then provide more detail on what government plans to do provided that this is principally consistent with and does not substantively change the model provided in the PM (though this does not eliminate the fact that it would have been better had the PM done so).

Beyond prompting requests for clarification from the Opposition, or criticism from them for the policy being vague and unclear, there is nothing more that should arise from a vague definition. Opposition might choose to argue that, given that the motion has been vaguely specified, a certain reasonable consequence or interpretation might be inferred from it. But they are not permitted to ignore the definition that was made, replace it with a preferred definition of their own choice, or claim that since they haven’t defined the motion clearly, OG are committed to defending very unreasonable applications of their policy.

To the extent that a Government team gains an advantage over another team because a previously vague policy has been later clarified or refined in a way that impairs their opponents’ ability to respond, that advantage should not be taken into account by the judges.

Worked Example: ‘**This House Would allow prisoners to vote**’

Example 1:

PM: ‘We define this motion as allowing prisoners the right to take part in elections.’

LO: ‘The PM has failed to confine this motion to adults in prison. Thus we must assume that children who are imprisoned will be allowed to vote, which is wrong as children are unfit to vote.’

DPM: ‘That’s clearly silly. Obviously child prisoners won’t be allowed to vote.’

The judge should conclude: The DPM is correct. The assumption made by the LO is unreasonable and must be rejected. OO receives no credit for their challenge.

Example 2:

PM: ‘We define this motion as allowing prisoners the right to take part in elections.’

LO: ‘The PM has failed to tell us which sorts of prisoners are allowed to vote. This definition is illegitimate because it doesn’t tell us which - and that might include murderers.’

DPM: ‘That’s silly! Of course our model doesn’t extend to murderers and the like, that would be completely unreasonable!’

The judge should conclude: Neither the DPM nor the LO are correct. There was nothing wrong with the PM’s definition, it merely left the opportunity for the Opposition teams to make arguments about why allowing murderers to vote would be a bad idea. It is not obvious that murderers were excluded from the PM’s definition, nor is it clear that they should be.

Challenging the Definition

If the definition provided by the OG is invalid, then it can be challenged. This must be done during the LO’s speech. As stated, the only grounds for claiming that a definition is invalid is if it meets one of the two squirrelling circumstances outlined above, or if it unfairly restricts the time and place of the debate. It is not enough for a definition to not seem ‘in the spirit of the motion’, or for a definition to have not been expected by other teams in the debate.

If a team challenges the definition, they must argue that the definition is illegitimate and explain why. In challenging the definition, the LO has two choices:

- a) Firstly, they can complain about the motion having been defined in an invalid way but proceed to debate it anyway. This is preferable if the motion proposed is not a fair reading of the motion but is still debatable. The debate then proceeds and is judged as per normal.
- b) Secondly, they can challenge the definition and redefine it. They should tell the judge and the other debaters what a proper definition would be and should then proceed to argue against that case. Where a team takes this option, it is advisable in some cases (though not required) for them to present ‘even-if’ analysis engaging with the OG’s definition of the motion and the material that stems from that definition, as well as their own.

Judges should not punish teams just for having a ‘definitional debate’. However, if teams engage in unnecessary definitional debates over reasonable definitions, this should be treated as self-penalising as they are wasting time on unpersuasive material at the cost of relevant arguments.

In **extremely rare cases**, OG may propose a wholly undebatable definition. If the LO neglects to challenge the definition, the MO may challenge this definition. In these scenarios, it is advisable for CO to offer Points of Clarification to OG. **These scenarios are exceedingly rare, and teams should be aware that attempting to challenge the definition when the motion is not wholly undebatable is likely to harm them. Teams should not pursue this strategy lightly.**

If the definition is challenged, judges must weigh the contributions teams made to the debate based on the accepted definition at the time they gave their speeches. To illustrate this, consider the following scenario. OG and OO agree on a definition, and OO clearly wins the top half debate based on this accepted definition. CG and CO then agree to expand the definition of the debate, and make contributions to the debate based upon the expanded definition. **Judges cannot then disregard OO because ‘the debate became about something else’ - rather, they must compare the relative contributions that each team made to the round, and consider moments where front half teams engage with back half teams and vice versa.**

Please bear in mind that definitional challenges are incredibly rare and more a ‘last resort’ than a first-line of defence against a Government case. Where a definition falls within one of the circumstances outlined above, it is often still advisable for a team to debate the motion as it has been defined, and avoid the procedural complexity of a definitional challenge taking away from their time to present substantive arguments.

Note that a definition cannot be attacked merely for being ‘the status quo’. Most motions will ask Government teams to defend the implementation of some sort of policy, which is likely to involve changing the world from the way it is at present. As such, if OG actually propose something which is identical to the status quo, this might be symptomatic of them failing to define the motion properly.

But as WUDC is an international tournament, with motions presumed to apply to many different countries which each have different existing policies, the mere fact that a definition is ‘status quo’ in some context is not a problem with the definition. For example, if the motion is **‘This House Would only have unicameral (single-chamber) legislatures’**, and OG propose that all democracies should have a single chamber parliament elected through a mix of constituency representatives and proportionate party-list members, they have proposed a policy which is the status quo in New Zealand. However, this would be a radical change for many democracies. Defining a debate in a way that happens to be status quo somewhere is not in and of itself a problem.

Whether a definition is valid or not, it is not the job of the judge to attack the definition, and judges should only worry about the definition if teams in the debate do. If the definition is successfully attacked as being vague, OG should be penalised only to the extent to which a lack of detail prevents teams from making arguments. Other teams should not be penalised for OG’s vagueness: judges should allow other teams to advance fair and reasonable assumptions, so long as they reasonably and logically follow from OG’s vague definitions.

Worked Examples: ‘This House Would allow prisoners to vote’

Example One:

PM: ‘We will allow all prisoners to vote in elections who have less than one week remaining in

their prison sentence.'

LO: 'This is clearly unfair as a definition of the motion as it unduly narrows the scope of the debate, but we'll oppose it anyway.'

The judge should conclude: The LO has made a correct challenge to the motion and the PM should be penalised.

Example Two:

PM: 'We would allow all wrongfully-accused prisoners to vote, having released them from prison.'

LO: 'This is a completely unacceptable narrowing and twisting of the definition to the point where Government teams have not argued that real prisoners should be allowed to vote. Since what they need to prove is that actual prisoners should be allowed to vote, that is what we will be arguing against. We oppose such a policy for the following reasons...'

The judge should conclude: The LO has done the right thing by replacing the unworkable definition with a workable one. Teams should follow the LO's lead and debate the motion as they have set forth.

2.10 Opposing the Debate

Government teams argue in favour of what the motion requires them to do or say is true. In a debate about a policy, the Opposition must say that we shouldn't do it; that is, that something is better than doing this policy. As with definitions of the debate by OG, the position Opposition chooses to defend can be the status quo in some countries, it can be something which is currently done nowhere, or it may be described as 'doing nothing' rather than 'doing the policy' (though naturally, teams doing this don't necessarily recommend wholesale government inaction, but are running the comparative line that 'whatever other broadly sensible relevant policies one is carrying out, the addition of this one makes things worse').

Counter-propping

So long as Opposition teams provide reasons not to do the policy, this is fine. It is not the Opposition team's burden to commit themselves to a particular or specific alternative course of action to the Government policy. However, they may choose to advance a 'counter-proposition' (or 'counter-prop' or 'countermodel'): this refers to a specific policy, or model, promoted by the OO (see section 2.7 on Motion Types for more information about policy debates). This course of action should only be undertaken when the motion type allows for a policy debate.

Just as only the OG has the right to set out a model for the Government side and must do so in the PM's speech, only the LO may set out a counter-proposition for the Opposition side. When advancing a counter-proposition, OO teams enjoy the same level of fiat as OG. The policy OO forwards as a counter-prop must use the same level of resources or fewer resources as the policy advocated for by OG. These resources could include the expenditure required to implement the policy, the use of some other resources that are not directly obtainable through expenditure, or the public will required to implement the policy.

Opposition counter-models, like OG models, can use different types of resources. For instance, a model may require public expenditure, public will, or the use of other resources. Opposition counter-props must not only use the same or fewer resources on net, but must also use the same or less of each type of resource.

Whether a resource is of the same ‘type’ as another is defined by whether the two resources could be converted into each other easily in the vast majority of circumstances. A resource used by OG cannot be converted into something different to be used in OO’s countermodel unless that conversion can easily happen in the vast majority of circumstances. For example, if a counter-prop were to cost far more money than OG’s model, that would likely be impermissible even if it was more popular and therefore used less ‘public will’. That is because it is not clear that ‘public will’ is obviously and easily convertible into money in the vast majority of instances.

For example, on the following motions:

Motion	OO Strategy	Evaluation
‘THW legalise marijuana’	‘We will take all of the public will saved by not legalising marijuana and will set up new clinics to deal with drug addiction’	Likely impermissible, as it is not clear that ‘public will’ is convertible to money in the vast majority of instances.
‘THW hire significantly more police officers’	‘We will instead redirect all of the money OG would use on police officers to hire social workers instead.’	Permissible counterprop (as long as the amount of money saved on police would be sufficient to cover the costs to hire social workers) as it is clear and obvious how an actor can spend less money on hiring police officers and more on social workers.

The counter-proposition proposed by the LO must be mutually exclusive with the policy mandated by the motion. It is important to note that a counter-proposition alters the comparative in the debate, as all teams need to compare the policy proposed by the Government with the counter-proposition rather than with the status quo. The debate is judged as per normal: teams advance arguments about the benefits and harms of both proposed models.

For example, on the motion ‘THW significantly increase taxes for individuals in the highest

income bracket', OO may counter-prop abolishing income taxes instead - a policy which is mutually exclusive to OG's model. It would, however, not be a counter-prop for OO to claim that they would educate individuals about philanthropy, as this is not mutually exclusive to the OG model.

Counter-props must also be related to one of the same broad policy areas as the OG model. While policy areas are not perfectly defined, judges can make reasonable assessments that policies fit into particular categories (e.g. education, health). Given that these areas are not perfectly delineated from one another, judges should not be overly restrictive in how they interpret what counts as being 'in the same broad policy area'. As a general rule, where the OO team proposes a reasonable countermodel which is intuitively related to the motion, it should be permitted.

If the original policy was aimed at multiple policy areas to a substantial degree, the counter model need only target one of them. For instance, in the motion '**This House Would subsidise youth football teams**', OO could countermodel that they would instead subsidise senior football teams, or instead subsidise art programmes for youths. They cannot, however, counter-model that they would instead buy spy satellites to combat terrorism.

OO does not need to advance a counter-proposition, and can still win the debate by arguing against the model proposed by the government (e.g., by arguing that OG's model will make the problem so much worse that inactivity is preferable or showing that OG's action will create a different, even larger problem). OO may also make arguments that suggest a range of viable alternative things would occur as a result of not implementing the motion. These are colloquially known as 'alternatives' and are to be evaluated in the same manner as any other argument (e.g. because teams are not using counter fiat, they are required to prove that the alternative is likely to occur). **This is not the same as advancing a counter-prop.**

For example, in the motion '**THW ban the sale of junk food**':

OO Strategy	Evaluation
'We will put a tax on junk foods'	Mutually exclusive counter-prop that alters the comparative of the debate.
'We will advertise promotions for healthy foods'	Not mutually exclusive to the OG model; not a counter-prop.
'We will use all of the money that would be spent on enforcing the ban to grow our nuclear arsenal to increase the nuclear deterrent'	Not related to the same broad policy area as the OG model; not a permissible counterprop.
'The effort needed for this is better used in other areas such as stopping smoking'	Not a counter-prop and thus no fiat power; OO should explain why these are mutually exclusive and why efforts are likely to be directed to

	other areas.
--	--------------

Worked Examples: ‘This House Would invade Yemen’

Example One:

PM: ‘We believe that the US should invade Houthi-controlled areas of Yemen at once and install a new government.’

LQ: ‘We believe that the US should invade Yemen at once, but they should also give economic assistance to a new Yemeni regime.’

The judge should conclude: OO’s counter-prop is not mutually exclusive with OG’s, and indeed accepts the premise of the OG’s case. OO is not actually opposing the motion.

Example Two:

PM: ‘We believe that the US should invade Houthi-controlled areas of Yemen at once and install a new government.’

LQ: ‘Rather than invading, the US should give military aid to rebel groups within Houthi-controlled areas of Yemen.’

The judge should conclude: OO’s ‘counter-prop’ is not strictly mutually exclusive with the OG’s case, but they have set it up as an alternative (in effect saying that ‘we suggest the model of a) not invading and b) giving military aid’). Depending on the arguments that follow, they may be able to successfully show that their model is preferable to OG’s, though it is valid for OG to accept the alternative argument as part of their own model.

2.11 Member Speeches – Extending the Debate

The MG and MO are each responsible for ‘extending’ the debate. An extension is defined as anything that hasn’t yet been said by that side of the debate. An extension can take a number of forms including:

- new arguments which have not yet been made in the debate
- new rebuttals to material raised by the other side
- new examples or case studies
- new characterisations
- new framing
- new analysis or explanation of existing arguments (including, but not limited to: new mechanisms, new characterisations etc)
- new applications of existing argumentation (e.g., if the Member points out that one of their opening half’s arguments is able to defeat a new argument from the other side), and
- new criteria for judging the debate or a new defence of existing criteria for weighing arguments.

In short, saying almost anything other than a word-for-word repetition of first-half's material will in some sense constitute an extension. In that sense, role fulfilment here is fairly easy and most extension speakers will succeed in fulfilling the bare minimum requirements of their role. There should be almost no instances of a team on Closing half adding no new material whatsoever.

However, a Closing team can only be credited for contributions to the debate that go beyond what has already been contributed by their opening half. Closing teams do not win through minimal additions to already well-substantiated points, but to the extent to which their contribution (including the summary) is meaningfully better than what has come before. A Closing team that contributes only the most minimal of extensions is unlikely to have contributed more persuasive material than their Opening. As a result, Closing teams do not defeat their opening half team merely by 'having an extension' (any more than OG teams win the debate for 'having a model'). A winning extension will bring out material that is most able to persuade the judge that the motion should be affirmed or rejected.

When judging the cases of Closing teams, judges should identify what is exclusively new coming from the Closing case, and then compare only exclusively new material to the Opening case (or to any other team in the debate).

Knifing

Closing teams should be consistent with their opening teams (for further clarification refer to 2.3 under the section 'Contradictions').

There are some rare exceptions, in which Closing teams do not have to be consistent:

- a) The Opening team has conceded the debate, or made an extremely damaging concession that makes the debate impossible to win from their side.
- b) The Opening team makes a concession that does not advance their own case, with the sole purpose of strategically narrowing their own Closing team's available arguments (for example, making an unsubstantiated claim in fifteen seconds that does not advance their case in any meaningful way and stating 'if CG tries to say the opposite of this claim, it is a knife').⁴
- c) OG has squirrelled the motion (or OO has made an invalid counter-prop).
- d) The opening team has made a clearly false factual statement that an ordinary intelligent voter would recognise as false (e.g., in a debate about space travel, claiming that the moon is made out of cheese).

To be clear, under these rare circumstances, Closing teams still have to be consistent with other things said by their opening - this is not a 'blank cheque' to ignore everything that an opening

⁴ This is distinct from the Opening team making a strategic concession in order to advance their own case, which is permissible. For example, on the motion '**This House Opposes Kamala Harris as the Democratic Presidential nominee in 2028**', if OG argues that Harris is more likely to lose than any other nominee, it is acceptable for OO to argue that Harris being the nominee is good because it makes it more likely that the Democrats lose the election, and this is good because the Democrats are bad - even though this involves conceding to OG that Harris is likely to lose and binds CO to agree. This is an acceptable strategy because OO is using this concession to advance their own case, not merely to bind their closing team to an unfavourable comparative.

team has said, just the parts that it would be implausible or unreasonable to expect a Closing team to defend. **To reiterate: these cases are extremely rare, and we would not expect teams to invoke these clauses.**

Furthermore, proposing a different metric by which the debate should be evaluated does not usually constitute a knife. For example, if OO claimed that the most important thing in the debate is human rights, it is permissible for CO to claim that geopolitical impacts are in fact more important.

2.12 Whip Speeches

Just like any other speaker, whip speakers may make new contributions to the debate, including material that (i) is responsive to claims made by other teams in the round (new rebuttal) or (ii) contributes additionally to material already present at member. However, Whip speakers may not introduce entirely new, non-responsive material that significantly alters the direction of the case as presented by their member speaker.

If a team does introduce such a contribution in the whip speech, judges should simply ignore it, and not afford it any credit. Adding material in this manner shouldn't be penalised beyond this - rather the judge removes the advantage afforded by the rule violation by ignoring the material presented.

When assessing whether new non-responsive material is creditable at whip, judges should consider whether the material **'significantly changes the direction of the case from member in a manner that other teams reasonably could not have predicted'**. Claims which do 'significantly change the direction of the case from member' and are non-responsive should not be credited.

The following is a non-exhaustive list of contributions which do not count as significantly changing the direction of the case from member, and are permissible for whips to engage in, even where they are non-responsive:

- new defences of arguments already made
- new explanations of previously-made arguments (including, but not limited to: new characterisation, new mechanisms etc)
- new examples to support existing arguments
- new explanation regarding the impact or prioritisation of existing lines of argumentation,
- new characterisation
- new framing
- new weighing to explain why their team's contributions are the most persuasive or important on their bench, or more important than the contributions on the opposing bench.

2.13 Equity

As well as following the rules of BP debating, Worlds also requires that all participants adhere to the tournament 'Equity Policy'. Judges have no authority to enforce the equity policy (but must obviously themselves follow it). Judges may not cut off a speaker for a perceived breach of equity except in the most extreme of situations, where an equity violation is severe enough to have

already disrupted the round and intervention is required to restore order.

Judges should not take the fact that they believe an equity violation has occurred into account when assessing who won a debate, or what speaker points to award. Judges are there to judge the debate, and should only penalise equity violations to the extent to which they make a speaker unpersuasive and/or are unfair on other teams or speakers. Judges cannot award a speaker zero speaker marks, or give their team an 'automatic fourth' on the basis of a breach of equity.

To resolve equity violations formally, debaters and/or judges should report them to the equity team who, in consultation with the CAP and the person making the complaint, will decide what course of action, if any, needs to be taken. However, being an objectionable speaker is generally not persuasive to the ordinary intelligent voter. A speaker who engages in, for example, racist behaviour is likely to be rendered less persuasive overall as a result of that material.

3. Additional Notes for Judges

Most of the information on how to judge debates and determine results was provided in Chapter 2 - as such **all judges must read Chapter 2 of this manual for guidance on judging**. This section simply focuses on a few additional issues of a largely administrative nature for judges: such as how to actually engage in the judging deliberation, fill in the ballot, deliver feedback to the debaters, and so forth.

3.1 Deciding the Results

Once the debate has finished, the debaters should leave the debate room, and the judges should collectively rank the four teams in order: first, second, third and fourth. Judges do this through a discussion (or 'deliberation') aimed at consensus - they do not simply each make up their minds and then vote, or engage in a battle with each other to 'win' the discussion. Judging panels are a team, and all members of the panel should view themselves as such - their job is to cooperatively decide on the best way to rank the four teams in the debate. Debates cannot result in a draw: one team must take the 'first', one team the 'second', one team the 'third', and one team the 'fourth'.

To repeat the core BP debating criterion on winning debates: judges assess which teams were most *persuasive* with respect to the *burdens* their side of the debate is attempting to prove, within the *constraints* set by the rules of BP debating. Judges should determine which team did the best to persuade them, by reasoned argument, that the motion ought to be adopted or rejected. The judges do so as the ordinary intelligent voter within the meaning outlined in section 2.2, and their assessments are always *holistic* and *comparative*: considering all the contributions each team made to the debate in aggregate, and comparing these to other teams. Teams cannot win or lose debates for isolated things they did, like setting up the debate well or contradicting another team on their side.

Crucially, there are no such things as 'automatic fourths' or 'automatic firsts'. This is a matter of logical necessity: however good or bad something a team does is, another team could always do exactly the same good or bad thing *and* do something else that made them even better or even worse.

Judges can and must assess how well-substantiated arguments are. This will inevitably involve some assessment of the quality of the supporting reasons offered for arguments; and, as noted in section 2, seriously implausible claims may constitute weak support for an argument in the eyes of the judges. But judges must exercise the minimum of personal evaluation in making such claims, and even seriously implausible arguments cannot be disregarded entirely by the judge if they haven't been rebutted - though they may have little persuasive value.

In an ideal world, teams will engage in extensive responses to each other's well-detailed points. In most of the debates that occur in the actual world, teams will often talk past each other and leave each other's points unchallenged. Under those circumstances, the judge will have to assess not only which arguments are most important, but equally which are most clearly proven. Unrebutted points that require the judge to make some logical leaps are often more persuasive than thoroughly-rebutted points and are always more persuasive than no points at all, but are not

preferable to a well-reasoned argument which rests on fewer unsubstantiated assumptions. What is and is not rebutted is therefore of vital importance to judging debates.

It is also important to identify correctly the direct engagement between specific teams. Just as OG cannot defeat the OO due to constructive arguments that CG provided, similarly, OG cannot defeat OO due to a rebuttal provided by CG. When comparing specific teams, we must take into account what those teams specifically engaged with, and had the opportunity to engage with.

Note that speakers don't have to use the word 'rebuttal' to respond to an argument. It may be tidier if they do, but judges should not ignore material that adequately deals with an argument just because the speaker doesn't point out that it does. Equally, this doesn't mean speakers should be 'punished' for not refuting everything: some claims do not do any harm at all to the opposite side. For example, in a debate about the legalisation of drugs, if the PM says 'pink elephants are cute because they have those nice ears and are a pleasant colour', this flawed argument can be safely left unrebutted as it isn't a reason to legalise drugs. There is, therefore, no need to point out that blue elephants are obviously more tasteful. So too, if they said 'some drugs are less harmful than others', this could also be ignored. While it is clearly more related to the debate than the cute pink elephants argument, it is pre-argumentative - that is, it has not yet been given sufficient surrounding words to actually provide a reason to do or not do the policy. The other side can quite happily say 'yes, some drugs are more harmful than others' and move on, or just ignore this argumentative non sequitur. Often as a judge, it can be tempting to complete arguments for teams that are interesting but pre-argumentative. Don't.

3.2 Judging Panels

Each judging panel will comprise a single 'Chair' and a number of additional judges termed 'Wings' (or 'Panellists'). It is the responsibility of the Chair to manage the deliberation between the judges in a manner that allows all judges to participate fully in the discussion, and produces a consensus decision and completed results sheet (known as a 'ballot') within the deliberation time limit: 20 minutes at this Worlds. Chairs of panels must manage their time accordingly, and recognise that the rules require a vote if no consensus has been reached early enough for the adjudication to complete in 20 minutes. Taking into account the time taken to decide on individual speaker points, this means you should consider a vote around 17 minutes into a discussion.

The opinions of Wings count just as much as the opinion of the Chair: the main difference is simply that Wings are just not tasked with chairing (i.e., managing) the discussion. Wings should treat the Chair with respect, and not interrupt/speak over them. If wings feel they are not being allowed to meaningfully participate in the discussion, or have concerns about the way in which they were treated by chairs, they should report this to the CAs via the judge feedback form, or to the Equity Officers (if necessary). They should, however, also be aware that Chairs are constrained by the time limit, and so may not be able to allot them as much time to speak as they might like. In return, Chairs should respect the opinions of Wings and give them sufficient opportunity to contribute to the discussion.

Voting on Pairwise Comparisons

At the end of the deliberation, if the voting members of the panel are not unanimous on the full order of teams, the chair must call a vote on all pairwise comparisons. The result of each contested comparison will be determined by majority vote, with the chair casting a tie-breaking vote in case the panel is tied. This is known as a split decision.

In case of a split decision, the chair may not assert that one team beats another ‘by transitivity’; they must call a vote on every split comparison. For example, take the case where the chair and panellist A would give 1st to OG; 2nd to OO; 3rd to CG; 4th to CO. Panellist B would give the 1st to OO; the 2nd to CG; the 3rd to OG; 4th to CO:

1. The panel unanimously agrees that OO beats CG, and that CO loses to each other team. The chair need not call a vote on any of these four comparisons.
2. The chair first calls a vote on OG against OO. OG beats OO on a 2-1 split.

Even though the majority of the panel has determined that OG beats OO, and all judges agree that OO beats CG, **this does not entail that OG ‘automatically’ beats CG**. As panellist B has CG beating OG, the chair **must** call a vote on the government bench comparison.

3. The chair finally calls a vote on OG against CG. OG beats CG on a 2-1 split.

In this illustration, following these votes, the 1st goes to OG, the 2nd goes to OO, the 3rd goes to CG, and the 4th goes to CO. The comparisons on the top half and the government bench were decided on a 2-1 split vote. The chair must announce as much in the oral adjudication.

The chair has discretion on the order in which they call votes between teams. However, the chair is **required** to call a vote on every non-unanimous comparison at some point in the deliberation.

If the pairwise votes lead to a single, coherent result for the round, then the voting process stops here. This will be true in the vast majority of rounds.

Cyclic Ties

A minority of debates may lead to a situation where, through a series of split decisions, three (or four) teams appear to be beating each other in a way that makes it impossible to produce a single decision. The simplest example is this:

Example A – Three-person Panel

For the sake of brevity, the notation ‘OG > OO’ represents OG **beating** OO. Each judge has the following call at the end of deliberation⁵:

- Chair: **OG** > **OO** > **CG**
- Panellist A: **OO** > **CG** > **OG**
- Panellist B: **CG** > **OG** > **OO**

After pairwise voting, we are left with three split decisions:

- **OG** beats **OO**, 2-1 split decision [Chair & B in majority]
- **OO** beats **CG**, 2-1 split decision [Chair & A in majority]

⁵ (CO omitted for simplicity, and unanimously 4th.)

- **CG** beats **OG**, 2-1 split decision [A & B in majority]

At the end of Example A, we are in a position where it is impossible to determine a ranking between the three teams. This scenario is known as a **cyclic tie** or a 'triangle call'.

We again emphasise that the chair may not claim that one team beats another by 'transitivity', and may not rely on 'transitivity' to resolve a cyclic tie. If a cyclic tie exists, the chair must resolve it using the **Ranked Pairs method**, described below.

The Ranked Pairs Method

In brief, the Ranked Pairs method sorts the pairwise comparisons between teams by margin of votes (with the chair resolving ties), and then removes any comparisons which conflict with higher-ranked ones. In essence, the narrowest part of the voting 'triangle' is removed to produce a coherent result.

The chair should take the following steps:

1. Tally each pairwise comparison and record its margin (majority minus minority).
2. Sort comparisons from largest to smallest margin; the chair breaks any ties (using their judgment).
3. Lock comparisons in that order, skipping any that would create a cycle with already-locked pairs.
4. Continue locking/skipping until a full, coherent ranking of teams is produced.

Ranked Pairs Explained

For the purposes of clearly illustrating the Ranked Pairs method, we will be using a larger panel of 5 judges.

Example B – Five-person Panel

Calls at the end of deliberation are as follows:

- **Chair** **OO** > **CG** > **OG** > **CO**
- **Panellist A** **OG** > **OO** > **CG** > **CO**
- **Panellist B** **CG** > **OG** > **OO** > **CO**
- **Panellist C** **CG** > **OG** > **OO** > **CO**
- **Panellist D** **OG** > **OO** > **CG** > **CO**

After pairwise voting, we are left with these decisions:

- **OG vs OO** **OG** wins 4-1 (Chair is the lone minority vote)
- **OO vs CG** **OO** wins 3-2 (Panellist B, Panellist C are the minority votes)
- **CG vs OG** **CG** wins 3-2 (Panellist A, Panellist D are the minority votes)
- **CO** loses to all other teams unanimously.

The chair identifies that pairwise voting has yielded a cyclic result between OG, OO, and CG, and proceeds with the Ranked Pairs method.

Step 1: Tallying

Start by tallying the votes for each of the six pairwise comparisons, and calculating the margin for each comparison. The margin of a pairwise vote is the number of judges in the majority minus the number of judges in the minority.

In our example, the margins are these:

Comparison	Vote Count	Margin
OG vs OO	4-1	3 votes
OG vs CG	2-3	1 vote
OG vs CO	5-0	5 votes
OO vs CG	3-2	1 vote
OO vs CO	5-0	5 votes
CG vs CO	5-0	5 votes

Step 2: Sorting

The pairwise comparisons are then **sorted** by margin, with the highest-margin comparison ranking 1st.

Sorted Comparisons (with ties)		
Rank	Comparison	Margin
1=	OG > CO	5 votes
1=	OO > CO	5 votes
1=	CG > CO	5 votes
4	OG > OO	3 votes
5=	OG < CG	1 vote
5=	OO > CG	1 vote

Where two comparisons have the same margin as each other, the two comparisons are tied. The chair must use their discretion to decide which comparison to rank higher. The chair may decide based on which of the two comparisons they thought was clearest, for instance, or which team the chair had placing higher in their estimation of the round.

In our example, the chair starts by resolving the ties between the three pairwise comparisons involving CO. It should be noted that the result of this tiebreak does not ultimately change the call.

The chair must then tiebreak between the two comparisons tied for 5th place, with a margin of 1 vote: (i) [OG < CG] and (ii) [OO > CG].

- You will recall that the chair's own ranking of these three teams was **OO > CG > OG**. The chair might therefore prefer (ii) [**OO > CG**] on that basis, as they have **OO** winning the round.
- Alternatively, the chair might instead prefer (i) [**OG < CG**] if they believe that **CG** beat **OG** far more clearly than **OO** beat **CG**.
- Chairs should use their good sense to resolve ties in ways they can justify. As a chair, you should be able to explain to teams why you thought one comparison was clearer than another.

The comparisons have now been sorted:

Sorted Comparisons (Final)	
Rank	Comparison
1.	OG > CO
2.	OO > CO
3.	CG > CO
4.	OG > OO
5.	OO > CG
6.	OG < CG

Step 3: Locking

Starting from the top of the ranked list, start 'locking in' pairs. Lock each pair in turn, unless that pair would create a cyclic result based on the pairs you've already locked in (i.e., a 'triangle'). If it does, skip it. Repeat until a full result is produced.

For Example B, the locking in step looks like this:

R.	Comparison	Action	Result	Round Result
1.	OG > CO	<i>Lock in</i>	OG beats CO	OG > CO
2.	OO > CO	<i>Lock in</i>	OO beats CO	OG > CO; OO > CO
3.	CG > CO	<i>Lock in</i>	CG beats CO	OG > CO; OO > CO; CG > CO
4.	OG > OO	<i>Lock in</i>	OG beats OO	OG > OO > CO; CG > CO
5.	OO > CG	<i>Lock in</i>	OO beats CG	OG > OO > CG > CO
6.	OG > CG	Skip: cyclic with (4) & (5)	<i>No effect</i>	

The final result is therefore **OG** > **OO** > **CG** > **CO**. The Ranked Pairs method has given the panel a way of determining which pairwise comparison should be ignored to resolve the (initially) cyclic result.

Note that the chair's tiebreaking decision to rank [**OO** vs **CG**] over [**OG** vs **CG**] had an effect on the final result. Had the chair decided the other way, the call would have been **CG** > **OG** > **OO** > **CO**. The chair therefore possesses some discretion in resolving ties.

Trainee Judges

Some judges in the tournament may be designated as 'trainees'. Trainee judges function exactly like Wing judges in every respect *except* that they do not get a vote in the eventual determination of the round's results. Trainee judges do still get to participate in the deliberation, and should follow, make notes on, and declare their views/rankings of the debate. Chair judges should give them equal opportunity to voice their views and other judges should engage with them in discussion directly. But the trainee does not get a say in deciding on the ultimate results of the debate, nor are they allowed to cast a vote in the event that there is no consensus among the panel. Being designated a 'trainee' should not be read as indicating that the CAP thinks a judge is bad. More usually it reflects that either the judge has limited judging experience, or that the CAP lacks information on the judge.

Chair, Wing and Trainee designations may change over the course of the tournament as the CAP gains more information about the judge in question, either through feedback from teams and panellists or through judging with them.

3.3 Managing the Discussion

In close rounds, it is to be expected that the judges on the panel may have different views on the debate. Therefore, achieving consensus and filling in the results ballot in 20 minutes is a difficult task, requiring careful management by the Chair. Here we sketch some suggestions for how this could be managed. These are not strict requirements - it is up to the Chair to manage the discussion in an effective way.

It is reasonable to take a few minutes to organise notes and confirm opinions individually prior to starting discussion. The Chair should then ask each Wing to give either a full ranking of the four teams or, at least, some indication of which teams they considered better or worse than each other. If Wings do not yet have a complete ranking, they should feel free to provide more general intuitions (e.g., 'top-half', 'bottom half', 'Government bench', 'Opposition bench'). That said, it is important that comparisons between teams be 'pairwise'. That is, if two teams are being compared, the contributions of another team are not relevant in this comparison. For example, a strong CG team will not strengthen the position of OG.

This is not binding, it is a working hypothesis which will evolve as the discussion progresses. *Wings should not feel any pressure to agree with one another or the Chair in their initial call, as there is no negative consequence or inference for changing your call.*

The Chair should then assess the level of consensus which exists. There are many possible combinations, but thankfully a few scenarios crop up fairly often.

- a) Everyone has exactly the same rankings - have a brief discussion to ensure rankings are the same for similar reasons. Move on to scoring.
- b) Everyone has the same except 1 person - ask them to defend their position. Be specific, tailoring the requested defence to the difference between the minority and majority opinion. If it is a difference of one team, focus on that team, etc.
- c) There is similarity in rankings but also some crucial differences - You agree on where 1 team is ranked or some relative rankings - everyone agrees OG is better than CG) Begin by establishing which discussions need to happen (i.e., there is disagreement about whether OO beats OG). Begin by consolidating the consensus that exists, and use this as a platform to break deadlocks.
- d) Chaos - There is no similarity between the rankings. Guide a discussion of each team's arguments, or, depending on what makes sense to you and in context, of the clashes between particular pairs of teams. These debates often hinge on how one argument was evaluated, so your aim is to detect such differences in interpretation. The initial discussion is intended to inform each other of your perspectives and find some level of common understanding. If two judges believe different arguments are central, frame a discussion about their relative priority. Get each judge to explain their position, and attempt to establish a metric for the importance of arguments in the debate.

After this brief discussion, rank the teams and compare again. If you have achieved some overlap, move on to the suggestions under (c) above. Vote if necessary.

In all deliberations, *judges should not feel under any obligation to stick to their original call just because it was their initial view* - flexibility and open-mindedness in the discussion is crucial, and deliberations should always aim at consensus. Such consensus is not, however, an ideal that is to be placed above the right result.

As such, judges should not 'trade' results in order to each get their own views somewhat represented in the final ranking - this is likely to produce a result that is impossible to coherently justify. If a judge believes that a team placed first and the other judges disagree, the former judge should try to advance their reasons. All judges must be flexible and willing to be persuaded, but if they are not persuaded, they should stick with what they believe to be right.

Please note that whilst achieving a consensus is ideal, it is not always possible. Opinions may not change or the time it would take to change them is longer than the time allocated. A split may at some points be a more accurate evaluation of what happened in the debate. Do not make decisions based on untidy compromises, but do not fear to call a vote on issues. During feedback, we expect you to explain the decision to use votes to the debaters and how the outcome of these votes affected the final call.

3.4 Filling in the Ballot

Decide the ranking first, with no consideration of speaker marks until this has been established. This reflects the fact that teams win debates, not speakers, and they win based on their aggregate contribution. We are not evaluating our aesthetic appreciation of the speeches (or

proxy-marking ‘team balance’): we’re assessing the team’s aggregate contribution. Imbalance within a team should be reflected by giving the speakers different speaker marks.

Once a ranking has been decided upon, the Chair should lead the panel in filling in the ballot. This involves recording the rankings and assigning ‘speaker scores’ - a score, from 50-100, for each speaker in the debate. The speaker point scale, with guidelines on how to award speakers, is attached as an appendix to the end of this manual. There are a few important rules about awarding speaker scores:

- Speaker scores are allocated on a consensus basis. In the event that judges are not able to come to consensus on speaker scores, the chair may elect to call a vote on speaker scores in the same manner that they may elect to call a vote on pairwise comparisons, with the chair holding the tie-breaking vote. As with voting on results, voting eligibility rules apply (i.e., trainees and judges that missed some speeches in the round are not eligible to vote on speaks).

Speaker scores should reflect the majority decision of the judges, not be a compromise between various opinions (i.e., don’t say ‘we think OG wins, but we can make sure the speaks reflect your different view’). If the majority doesn’t think a relative ranking is close, there is no reason that the speaker scores suggest otherwise.

- The combined speaker scores for the two speakers’ on each team must be compatible with the ranking they received.

The team that placed first must have a higher combined speaker score than the team that placed second, the team that placed second must have a higher combined speaker score than the team that placed third, and so on. Teams cannot be given the same total speaker score - there must be at least a one point difference in the total speaker score of each team.

- Judges should assess all speakers in a fair manner and must take note of the fact that neither language proficiency nor accent influence a speaker’s speaker score.

3.5 Announcing the Result

The chair of the panel delivers the oral adjudication (OA). In the case that the chair loses a vote and feels unable to justify the call, they may retire from this position and ask one of the wing judges who voted in the majority to deliver all or part of the adjudication. If the chair chooses to give the adjudication, this must be to defend the majority position, although the chair should overtly state they disagreed with the majority.

The primary aim of an OA is to convey to the teams the reasoning of the panel in ranking the teams as they did. The OA should therefore present a logical argument for the ranking, using as evidence the arguments made in the debate and how they influenced the judges.

Leaking refers to any sharing of classified information by judges to anyone other than the adj core of the tournament. This information includes, but is not limited to: speaker scores before the release of the speaker tab, and the results of closed rounds. Judges who are caught having leaked information shall be severely punished.

3.6 Feedback on Adjudicators

Teams and judges are required to submit feedback on one another. There are three types of feedback:

- teams' feedback on the judge who delivered the adjudication
- chairs' feedback on wings and trainees
- wings' and trainees' feedback on chairs.

In scenarios in which multiple people (for instance, the chair and a wing in a split decision) contributed to the OA, teams may give multiple pieces of feedback, one for each contributor. Each type is important. The only way CAPs can effectively assess and allocate judges is if everyone participates in providing feedback.

Appendix A: The WUDC Speaker Scale⁶

The mark bands below are rough and general descriptions; **speeches need not have every feature described to fit in a particular band**. Many speakers will range across multiple bands depending on the feature assessed - for example, their style might appear in the 73-75 range, while their engagement might be closer to the 67-69 bracket, and their argumentation closest to the 70-72 range. Judges should not treat any individual feature as decisive in and of itself, but should rather aim to balance all features of the speech to come to the speaker score that seems most appropriate. Throughout this scale, **‘arguments’ refers both to constructive material and responses**. Judges should assess all speakers in a fair manner and must take note of the fact that **neither language proficiency nor accent influence a speaker’s speaker score**. Please use the full range of the scale.⁷

Score	Qualitative Comments
95-100	• Plausibly one of the best debating speeches ever given; • It is incredibly difficult to think up satisfactory responses to any of the arguments made; • Flawless and compelling arguments.
92-94	• An incredible speech, undoubtedly one of the best at the competition; • Successfully engaging with the core issues of the debate, arguments exceptionally well made, and it would take a brilliant set of responses to defeat the arguments; • There are no flaws of any significance.
89-91	• Brilliant arguments successfully engage with the main issues in the round; • Arguments are very well-explained and illustrated, and demand extremely sophisticated responses in order to be defeated; • Only very minor problems, if any, but they do not affect the strength of the claims made.
86-88	• Arguments engage with core issues of the debate, and are highly compelling; • No logical gaps, and sophisticated responses required to defeat the arguments; • Only minor flaws in arguments.
83-85	• Arguments address the core issues of the debate; • Arguments have strong explanations, which demand a strong response from other speakers in order to defeat the arguments; • May occasionally fail to fully respond to very well-made arguments; but flaws in the speech are limited.

⁶ Speaker scale initially created by Sam Block, Jonathan Leader Maynard and Alex Worsnip and updated by the Warsaw EUDC CAP.
⁷ See section 3.4 for more detailed information about filling in the ballot and determining speaker scores.

Score	Qualitative Comments
79-82	- Arguments are relevant, and address the core issues in the debate; - Arguments well made without obvious logical gaps, and are all well explained; - May be vulnerable to good responses.
76-78	- Arguments are almost exclusively relevant, and address most of the core issues; - Occasionally, but not often, arguments may slip into: (i) deficits in explanation, (ii) simplistic argumentation vulnerable to competent responses or (iii) peripheral or irrelevant arguments; - Clear to follow, and thus credit.
73-75	- Arguments are almost exclusively relevant, although may fail to address one or more core issues sufficiently; - Arguments are logical, but tend to be simplistic and vulnerable to competent responses; - Clear enough to follow, and thus credit.
70-72	- Arguments are frequently relevant; - Arguments have some explanation, but there are regular significant logical gaps; - Sometimes difficult to follow, and thus credit fully.
67-69	- Arguments are generally relevant; - Arguments almost all have explanations, but almost all have significant logical gaps; - Sometimes clear, but generally difficult to follow and thus credit the speaker for their material.
64-66	- Some arguments made that are relevant; - Arguments generally have explanations, but have significant logical gaps; - Often unclear, which makes it hard to give the speech much credit.
61-63	- Some relevant claims, and most will be formulated as arguments; - Arguments have occasional explanations, but these have significant logical gaps; - Frequently unclear and confusing; which makes it hard to give the speech much credit.
58-60	- Claims are occasionally relevant; - Claims are not be formulated as arguments, but there may be some suggestion towards an explanation; - Hard to follow, which makes it hard to give the speech much credit.

Score	Qualitative Comments
55-57	• One or two marginally relevant claims; • Claims are not formulated as arguments, and are instead are just comments; • Hard to follow almost in its entirety, which makes it hard to give the speech much credit.
50-55	• Content is not relevant; • Content does not go beyond claims, and is both confusing and confused; • Very hard to follow in its entirety, which makes it hard to give the speech any credit.

Appendix B: Chair Feedback Scale⁸

The mark bands below are rough and general descriptions; judges need not satisfy every feature described to fit in a particular band.

Score	General Description	Qualitative Comments
10	Exceptional	Accuracy: Extremely accurate call, reflected through precise appreciation and very meticulous assessment of ‘close’ comparisons between teams; comprehensive recognition of all necessary inter-team comparisons. Reasoning/Justification: Extremely well-justified justification, evidenced by flawlessly or near-flawlessly outlined explanations that are in-depth, insightful, and nuanced; explicit identification and strong justification for any weighing metrics or assumptions employed in judging. Discussion: Offers highly astute and insightful comments on the debate; highly efficient, and demonstrates profound acumen in managing the panel discussion and (where appropriate) offering constructive feedback to teams.
9	Excellent	Accuracy: Very accurate call, reflected through appreciation and correct assessment of ‘close’ comparisons between teams; comprehensive recognition of most necessary inter-team comparisons. Reasoning/Justification: Very well-justified justification, evidenced by well-outlined explanations that are in-depth, insightful, and nuanced; good attempts made to justify weighing metrics in judging. Discussion: Offers very insightful comments on the debate; consistently efficient, and demonstrates effectiveness and judgement in managing the panel discussion.

⁸

Wing and Trainee scale originally created by the 2019 Athens EUDC CAP.

Score	General Description	Qualitative Comments
8	Very Good	Accuracy: Accurate call, reflected through largely correct judgement regarding ‘close’ comparisons between teams; detailed recognition of most necessary inter-team comparisons. Reasoning/Justification: Comprehensively justified justification, evidenced by well-outlined explanations that are in-depth and nuanced; very occasional slippage into minor assumptions or personal biases in judging, or minor lack of clarity in one or more inter-team comparisons; metrics for judging are identified but not explicitly justified. Discussion: Offers mostly insightful comments on the debate; largely efficient, and demonstrates effectiveness in managing the panel discussion.
7	Good	Accuracy: Accurate call, reflected through generally correct rankings but potentially wrong regarding ‘close’ comparisons between teams; careful acknowledgment of most necessary inter-team comparisons in consideration. Reasoning/Justification: Generally well-justified justification, evidenced by well-outlined explanations; occasional slippage into minor personal biases and assumptions, or minor lack of clarity in some inter-team comparisons. Discussion: Offers generally relevant comments on the debate; efficient with occasional slip-ups and flaws or imbalance in managing discussion; demonstrates an appropriate level of judgement (at times limited) in oral adjudication.
6	Above Average	Accuracy: Mostly accurate call, although may fail to get ‘close’ comparisons between teams correct. Reasoning/Justification: Good attempt at justifying decision; explanations demonstrating some appreciation of key clashes and how they are resolved; occasional slippage into minor or insignificant personal biases and assumptions; lack of clarity in some inter-team comparisons. Discussion: Offers some helpful or useful comments on the debate; somewhat inefficient and barely satisfactory at leading discussion; demonstrates a lack of understanding of the key issues in the debate in oral adjudication.

Score	General Description	Qualitative Comments
5	Average	Accuracy: Broadly accurate call that gets the ‘obvious’ clashes correct; may fail to produce accurate judgement regarding ‘close’ comparisons, or may neglect a significant but not substantial part of the debate. Reasoning/Justification: Some attempt at justifying decision; explanations demonstrating some appreciation of key clashes and issues; regular slippage into personal biases and assumptions, some of which may undermine the quality of the justification; lack of clarity regarding specific inter-team comparisons. Discussion: Mostly inefficient at leading discussion; at times, struggles with catering to one or more voices on panel without reason; demonstrates lack of mature judgement in oral adjudication.
4	Below Average	Accuracy: Inaccurate call that nonetheless identifies the ‘obvious’ rankings correctly; call reflects one or more misunderstandings of the debate; some inability to track important arguments/responses. Reasoning/Justification: Unsatisfactory attempt at justifying decision; explanations demonstrate some appreciation of key clashes and issues, but may not warrant or justify the posited call; frequent slippage into personal biases and assumptions, some undermining the quality of the justification; lack of clarity regarding most inter-team comparisons. Discussion: Incompetent at managing discussion; struggles to consider or include all members on panel; somewhat irrelevant in oral adjudication.
3	Poor	Accuracy: Inaccurate call failing to identify one or more of the ‘obvious’ rankings correctly; call reflects several misunderstandings of the debate, some of which may be fundamental; some inability to track important arguments/responses. Reasoning/Justification: Poor attempt at justifying decision; explanations demonstrating no appreciation of key clashes and issues; frequent slippage into personal biases and assumptions, most of which certainly undermine the quality of the justification and severely distort the results; lack of clarity regarding most inter-team comparisons; justification occasionally slips into utter irrelevance. Discussion: Incompetent at managing discussion; struggles to consider or include all members on panel; mostly irrelevant in oral adjudication.

Score	General Description	Qualitative Comments
2	Very Poor	Accuracy: Wildly inaccurate call that completely fails to identify more than one of the ‘obvious’ rankings correctly; call reflects several core misunderstandings of the debate; clear inability to track important arguments/responses. Reasoning/Justification: Little to no attempt at justifying decision; explanations demonstrating no appreciation of key clashes and issues; frequent slippage into personal biases, irrelevance and assumptions, that cumulatively undermine the quality of the justification and severely skew the results; lack of clarity regarding most inter-team comparisons Discussion: Very incompetent at managing discussion; struggles to consider any views of all members on panel; irrelevant and potentially counterproductive in oral adjudication.
1	Abysmal	Accuracy: Completely inaccurate call that absolutely fails to identify more than one of the ‘obvious’ rankings correctly; call reflects a fundamental and foundational misunderstandings of both the debate and British Parliamentary debating in general; clear inability to track important arguments/responses. Reasoning/Justification: Effectively no rationalisable attempt at justifying decision; explanations demonstrating no or deeply erroneous appreciation of key clashes and issues; consistent slippage into unwarranted personal biases and assumptions that cumulatively undermine the quality of the justification and severely skew the results; utter irrelevance. Discussion: Entirely incompetent at managing discussion; struggles to consider any views of all members on panel; irrelevant and very counterproductive in oral adjudication.

Appendix C: Panellist and Trainee Feedback Scale⁹

The mark bands below are rough and general descriptions; judges need not satisfy every feature described to fit in a particular band.

Score	General Description	Qualitative Comments
10	Exceptional	Accuracy: Extremely accurate call, reflected through precise appreciation and very meticulous assessment of ‘close’ comparisons between teams (reflected through speaker scores); comprehensive recognition of all necessary inter-team comparisons. Reasoning/Justification: Extremely well-justified justification, evidenced by flawlessly or near-flawlessly outlined explanations that are in-depth, insightful, and nuanced; explicit identification and strong justification for any weighing metrics or assumptions employed in judging; certainly should chair. Discussion: Outstanding contribution to the discussion that reflects exceptional judgement concerning what is relevant and useful to discussion, with a clear sense of prioritisation; highly helpful; incisive in commentary.
9	Excellent	Accuracy: Very accurate call, reflected through appreciation and correct assessment of ‘close’ comparisons between teams (reflected through speaker scores); comprehensive recognition of most necessary inter-team comparisons. Reasoning/Justification: Very well-justified justification, evidenced by well-outlined explanations that are in-depth, insightful, and nuanced; good attempts made to justify weighing metrics in judging; should chair. Discussion: Valuable contribution to the discussion that reflects good judgement concerning what is relevant and useful to discussion; very helpful.

⁹ Wing and Trainee scale originally created by the 2019 Athens EUDC CAP.

Score	General Description	Qualitative Comments
8	Very Good	Accuracy: Accurate call, reflected through largely correct judgement regarding ‘close’ comparisons between teams; detailed recognition of most necessary inter-team comparisons. Reasoning/Justification: Comprehensively justified justification, evidenced by well-outlined explanations that are in-depth and nuanced; very occasional slippage into minor assumptions or personal biases in judging, or minor lack of clarity in one or more inter-team comparisons; metrics for judging are identified but not explicitly justified; high potential to chair. Discussion: Comprehensive contribution to the discussion that reflects good judgement concerning what is relevant and useful to discussion; very helpful.
7	Good	Accuracy: Accurate call, reflected through generally correct rankings but potentially wrong regarding ‘close’ comparisons between teams; careful acknowledgment of most necessary inter-team comparisons in consideration. Reasoning/Justification: Generally well-justified justification, evidenced by well-outlined explanations; occasional slippage into minor personal biases and assumptions, or minor lack of clarity in some inter-team comparisons; has potential to chair. Discussion: Good contribution to the discussion that reflects mostly good judgement about what is relevant and useful to discussion; helpful, with only minor lapses in attention and judgement.
6	Above Average	Accuracy: Mostly accurate call, although may fail to get ‘close’ comparisons between teams correct. Reasoning/Justification: Good attempt at justifying decision; explanations demonstrating some appreciation of key clashes and how they are resolved; occasional slippage into minor or insignificant personal biases and assumptions; lack of clarity in some inter-team comparisons. Discussion: Good contribution to the discussion that reflects mostly good judgments concerning what is relevant to discussion; helpful, with some lapses in attention and judgement.

Score	General Description	Qualitative Comments
5	Average	Accuracy: Broadly accurate call that gets the ‘obvious’ clashes correct; may fail to produce accurate judgement regarding ‘close’ comparisons, or may neglect a significant but not substantial part of the debate. Reasoning/Justification: Some attempt at justifying decision; explanations demonstrating some appreciation of key clashes and issues; regular slippage into personal biases and assumptions, some of which may undermine the quality of the justification; lack of clarity regarding specific inter-team comparisons. Discussion: Average contribution to the discussion that reflects some judgement concerning what is relevant to discussion; mostly helpful, but may be unresponsive to prompts or generic at times.
4	Below Average	Accuracy: Inaccurate call that nonetheless identifies the ‘obvious’ rankings correctly; call reflects one or more misunderstandings of the debate; some inability to track important arguments/responses. Reasoning/Justification: Unsatisfactory attempt at justifying decision; explanations demonstrate some appreciation of key clashes and issues, but may not warrant or justify the posited call; frequent slippage into personal biases and assumptions, some undermining the quality of the justification; lack of clarity regarding most inter-team comparisons. Discussion: Average contribution to the discussion that can be at times irrelevant; sometimes helpful, but frequently unresponsive to prompts or generic.

Score	General Description	Qualitative Comments
3	Poor	Accuracy: Inaccurate call failing to identify one or more of the ‘obvious’ rankings correctly; call reflects several misunderstandings of the debate, some of which may be fundamental; some inability to track important arguments/responses. Reasoning/Justification: Poor attempt at justifying decision; explanations demonstrating no appreciation of key clashes and issues; frequent slippage into personal biases and assumptions, most of which certainly undermine the quality of the justification and severely distort the results; lack of clarity regarding most inter-team comparisons; justification occasionally slips into utter irrelevance. Discussion: Below-average contribution to the discussion that reflects somewhat flawed understanding; rarely helpful; generic or occasionally unhelpful commentary.
2	Very Poor	Accuracy: Wildly inaccurate call that completely fails to identify more than one of the ‘obvious’ rankings correctly; call reflects several core misunderstandings of the debate; clear inability to track important arguments/responses. Reasoning/Justification: Little to no attempt at justifying decision; explanations demonstrating no appreciation of key clashes and issues; frequent slippage into personal biases, irrelevance and assumptions, that cumulatively undermine the quality of the justification and severely skew the results; lack of clarity regarding most inter-team comparisons Discussion: Poor contribution to the discussion; unhelpful; at times counterproductive to discussion.

Score	General Description	Qualitative Comments
1	Abysmal	Accuracy: Completely inaccurate call that absolutely fails to identify more than one of the ‘obvious’ rankings correctly; call reflects a fundamental and foundational misunderstandings of both the debate and British Parliamentary debating in general; clear inability to track important arguments/responses. Reasoning/Justification: Effectively no rationalisable attempt at justifying decision; explanations demonstrating no or deeply erroneous appreciation of key clashes and issues; consistent slippage into unwarranted personal biases and assumptions that cumulatively undermine the quality of the justification and severely skew the results; utter irrelevance. Discussion: Very poor contribution to the discussion; highly obstructionist; detrimental to the panel.