

The Discussion Divide Playbook

Formats, processes, and GenAI prompts for better disagreement A companion to “Bridging the Discussion Divide with GenAI” — PerpetualInnovation.org

The goal is not agreement. The goal is a better disagreement. GenAI is the third participant, not the referee.

Start Here

Pick a path.

Want a ready-made program, already tested, often free? Go to the next section, pick one, and stop reading this. Want to run something yourself, today, at your own table? Go to The Six Moves. Want to work one issue end to end? Go to The Discussion Bridge.

Agree on a stop word first. Any participant, any time, no explanation required. The discussion pauses. Use “Pluto,” “Reset,” or “Process check.”

Pluto earns its place because of what that argument actually was. The object never changed — same mass, same orbit — when astronomers reclassified it. Everyone agreed about the thing and disagreed about the label. Discussions do this constantly. So when someone says “Pluto,” the question is: *are we disagreeing about reality, or about what to call it?* Sometimes the label matters. Sometimes it doesn’t. Knowing which saves twenty minutes.

Know when to stop entirely. This is a discussion resource — not counseling, therapy, legal advice, crisis intervention, or a substitute for professional mediation. Stop if anyone is afraid of another participant, can be punished afterward for what they say, or is in crisis.

You do not need to assess the situation correctly. You only need to stop. In the US: **988** for crisis, by call, text, or chat at 988lifeline.org. **1-800-799-7233** for domestic violence. Appendix A has the full limits — read it before facilitating for anyone other than yourself.

Use What Already Exists

A mature field has been doing this for decades. Several of these are free. If one fits your situation, use it instead of anything in this playbook.

Program	What it is	Time	Cost
Braver Angels Red/Blue Workshop	7 conservative-leaning and 7 progressive-leaning participants, trained moderator	~3 hours	Free or low
StoryCorps One Small Step	Two strangers of opposing views record a 50-minute conversation about their lives, not politics	~75 min	Free
Living Room Conversations	100+ free guides for about six people, self-run, no facilitator	~90 min	Free

National Issues Forums	Issue guides framing three or four approaches with tradeoffs, moderated deliberation	~2 hours	Low
Essential Partners	Reflective Structured Dialogue; free field guide and question-design guide	Varies	Free guides
Better Arguments Project	Aspen Institute; five principles plus free workbooks for workplaces and schools	Varies	Free
#ListenFirst Coalition	Directory of roughly 500 organizations — find a practitioner near you	—	Free

A note on One Small Step. If you only try one thing on this list, try that one. It is the closest existing program to the technique with the strongest research behind it, and StoryCorps has built a platform that lets anyone participate remotely.

So what is this playbook for? Three situations the programs above do not cover. Two people at a kitchen table who are never going to book a workshop. A board or team that has twenty minutes, not an afternoon. And the part none of them address — how to use a large language model as a third participant without letting it become the judge.

The Six Moves

Each move stands alone. Run one, or string several together. Most work without GenAI; where it appears, its job is to summarize, classify, question, and compare evidence. People decide what matters.

Each carries an evidence label, because a playbook about weighing evidence should say which of its own advice rests on any. **Supported** means direct experimental evidence. **Consistent** means it aligns with deliberation research but has not been tested as a discrete technique. **Framework** means it is offered as practice, not finding.

Effect sizes in this field are modest and durable — not transformative. Deep canvassing runs $d = 0.04$ to 0.08 . Braver Angels reports about an 8-point improvement in feeling toward the other side. Expect a clearer disagreement, not a changed mind.

1. Find the Shared Goal

10 minutes · Consistent

Start at the proposition both sides accept, so the disagreement narrows before anyone defends anything.

Run it. Each person states, in one sentence, the outcome they want. The other restates it neutrally and is corrected until it is fair. List what overlaps. Only then name where the approaches differ.

Example. A school board argument about a reading curriculum.

Opening: “You’re trying to indoctrinate children.” / “You’re trying to ban books.”

As outcomes: *“I want my kids reading things I can discuss with them at home.” / “I want my kids to encounter people whose lives are different from ours.”*

Both want children reading demanding material. Both want a say. The disagreement is about who decides which books, and by what process.

When it goes wrong. Someone states a method as an outcome — “I want the policy repealed” is a method. Ask what the repeal would protect.

Debrief. Did the disagreement narrow?

2. The Fair Summary

10–15 minutes · Supported

Replace the version of the other side you have been arguing with, with the one they recognize.

The research here is counterintuitive. Across 25 experiments, Eyal, Steffel and Epley found that *imagining* another person’s perspective did not improve accuracy — and tended to reduce accuracy while raising confidence. Asking works. Imagining does not. So the value is entirely in the correction loop. A summary that needs no correction usually means nobody checked.

Run it. Person A summarizes Person B’s position. B corrects it, as many rounds as needed, until B says it is fair. Reverse. Each then names one legitimate concern in the other’s position.

Example. Two neighbors on documentary proof of citizenship for voter registration.

“You think we should just let anyone vote.” → “I think noncitizen voting is rare enough that a documentation requirement will block more eligible citizens than ineligible ones.”

“You think people should have to jump through hoops.” → “I think verification has value even when the detected rate is low, and I’d accept a one-time burden to get it.”

Neither moved. Both are now arguing with someone who exists.

When it goes wrong. The summary is accurate but loaded — *“You believe fraud doesn’t matter”* — and gets rejected while the summarizer grows frustrated. The rule is unconditional: they decide when it is fair. After three failed rounds, switch to the Two-Laptop variant.

Debrief. Was the corrected version harder to argue with?

Prompt. *Using only what each participant actually said, write a neutral summary of each position. Do not infer motives. Flag any place where one person characterized the other rather than stating their own view.*

3. One Real Story

20–30 minutes · Supported

Each person tells one specific true story about how this issue touched a real life. The listener asks questions and does not rebut.

This is the move most likely to be skipped and the one with the strongest evidence. Kalla and Broockman ran three field experiments in which conversations built on argument alone produced **no measurable effect**. The same conversations, plus a non-judgmental exchange of personal narrative, durably reduced exclusionary attitudes for at least four months.

The active ingredient is the non-judgment. If a story is heard as material to be used later, the effect disappears.

Run it. One person tells a specific, concrete story — a particular person, a particular week, not an illustration of a principle. The listener may ask questions only to understand, and may not respond to the substance. Reverse. Return to analysis afterward.

Example. A nonprofit board deciding whether to means-test a service.

Not a story: *“Means testing creates barriers for the people who need help most.”*

A story: *“In March a woman came in with her daughter’s immunization records in a grocery bag, because that was the only place she had to keep paper. It took four visits to assemble what we asked for. On the fourth she said she’d already been turned away somewhere else for the same reason and had almost not come back.”*

The second can be asked about. The first can only be agreed with or not.

When it goes wrong. The story is an argument with characters in it — redirect to specifics: when, who, what happened next. The listener rebuts — stop and restate the rule. Or someone offers more than they are ready to have discussed. **Anyone May Pass matters more here than anywhere else. Say so out loud before starting.**

Do not run this where a participant would be describing their own trauma to someone who disputes it. If the topic is that close, use StoryCorps One Small Step instead — it is built for exactly this and it is free.

Debrief. What did you learn that an argument could not have told you?

4. Sort the Statements

15 minutes · Consistent

Separate the kinds of claims that get argued as if they were the same kind.

Run it. Collect the strongest statements made so far. Sort each: verifiable fact, assumption, forecast, value judgment, or concern. Mark which can be checked. Discuss one category at a time.

Example. A workplace disagreement about returning to the office.

“Productivity dropped 12% after we went remote” — fact, checkable “People are less committed when they’re not in the building” — assumption “We’ll lose our culture within two years” — forecast “People should be trusted to manage their own time” — value “I’m worried I’ll be passed over if I’m not visible” — concern

Only the first can be settled by looking something up. The argument had treated all five as the same kind of claim.

When it goes wrong. The group argues about the sorting. Productive for about five minutes, not after. If a statement won’t sort cleanly, it is usually two statements — split it.

Debrief. Which disagreements were factual, and which were about values or acceptable risk?

Prompt. *Classify these statements as verifiable facts, assumptions, forecasts, value judgments, or concerns. Explain any uncertain classification. Do not treat an emotional stake as evidence for or against a factual claim.*

5. Test the Claims

20–25 minutes · Mixed — see caution

Two tests: what would change your mind, and what is the claim actually resting on.

Run it. Each person writes one piece of evidence that would weaken their position, and one observable outcome that would show their preferred approach is not working. Compare. Then take the two or three central factual claims and find the best available source for each — ranking roughly from primary documents and government data, through peer-reviewed research and transparent research organizations, to journalism that shows its evidence, to advocacy sources, which are useful for understanding positions but not as sole authority on disputed facts. Mark what stays uncertain.

Caution. The first half descends from research finding that asking people to explain *how* a policy works reduces false confidence and attitude extremity. That finding has three preregistered replication failures, and a later study found it holds for consequentialist disputes but not for issues grounded in protected or sacred values.

Expect it to work on arguments about means and outcomes, and to fail on arguments rooted in identity or moral absolutes. **When it fails, that is information, not a facilitation error.** If no evidence could matter to either side, the group has just learned it is in a values dispute — which is what it needed to know.

Example.

“If a state adopted this and registration among eligible seniors didn’t measurably drop, I’d have to rethink the cost side.”

“Honestly, nothing about the numbers. I think it’s wrong in principle.”

The second answer is more useful. It ends the search for a decisive study.

When it goes wrong. Someone treats the question as a trap. Rephrase away from the person: not *what would change your mind*, but *what would this look like if it were not working?*

Debrief. Were the tests realistic and observable? Did a strongly held claim turn out to rest on weak evidence?

Prompt. *For each central claim, identify the strongest available source category and explain how directly it supports the claim. Separate primary evidence from interpretation and advocacy. Weight claims by evidence quality, not repetition. Identify what should be independently verified.*

6. Weigh the Tradeoffs

20–25 minutes · Consistent

Move past for-and-against by naming what each option gains, risks, and costs — and who carries each.

Run it. List the real options, including doing nothing. For each: likely benefits, likely costs, who bears the burden, what stays uncertain. Mark which tradeoffs are measurable and which need a value judgment. Then name the question the discussion has been avoiding.

Example. A town considering a downtown parking requirement.

The argument had been “pro-business” versus “anti-car.” The table produced something else: benefits fall on existing businesses, costs fall on future ones, and nobody in the room represented future ones. That was the real disagreement and it had never been said aloud.

When it goes wrong. It becomes a scoring exercise and someone starts counting rows. It is not a vote. Ask which single row would change their mind if it moved.

Debrief. Which tradeoff is the actual source of the disagreement? Who bears it?

Prompt. *Build a neutral tradeoff table for these options, including doing nothing. Include benefits, costs, burden distribution, evidence strength, and major uncertainties. Do not recommend a winner.*

Two Full Processes

The Discussion Bridge

30–45 minutes · Supported

1. Each person explains their position in their own words.
2. A facilitator or GenAI summarizes both without evaluating them. Participants correct until each says theirs is fair. **Do not proceed until they do.**
3. Separate factual claims, assumptions, forecasts, values, and concerns.

4. Fact-check the central claims and name what stays uncertain.
5. The people — not the AI — weigh the tradeoffs and decide what, if anything, follows.

Structured deliberation of this kind has measurable effects. Stanford's America in One Room found significant depolarization across immigration, the economy, healthcare, foreign policy, and the environment, with parties moving closer on 22 of 26 proposals.

When it goes wrong. Groups rush step 2 because it feels like preamble. It is the step with the most evidence behind it. Short on time? Cut step 4, keep step 2.

Variant: Two Laptops

When conversation keeps collapsing into characterizing the other's motives, remove the characterization entirely.

Each participant writes their own position independently. Both submit to the same GenAI conversation. GenAI drafts a neutral summary of each. **Neither may characterize the other's position at any point** — each may only correct the summary of their own. Then ask GenAI to map agreement, disagreements evidence could resolve, and differences in values or risk tolerance. Pick one resolvable question to follow up.

Use when the Fair Summary has failed twice.

Coming Back

20 minutes, later · Framework

A difficult discussion is not something you finish. Record what was learned, what changed, what stayed unresolved, and what evidence would help next. Then actually return with it and update the shared fact base.

This is the Perpetual Innovation™ pattern applied to a conversation: assess, improve the questions, examine the evidence, act on human judgment, regenerate as new information arrives.

Prompt. *Summarize what changed in this discussion: new facts learned, assumptions revised, questions still open, value differences remaining, and the best evidence to gather before the next round.*

Facilitator Notes

Start with fairness, not persuasion. People examine evidence after they believe they were represented accurately, not before.

Emotion explains interpretation. It does not prove claims.

Do not manufacture false balance. Same standard for every claim is not the same as the same outcome for every claim.

Keep sources visible. Generative systems can cite a real source that does not support the claim attached to it — NIST calls this confabulation.

Agreement is optional. A narrower, clearer, researchable disagreement is success.

Human judgment stays decisive.

Appendix A — Safety, Limits, and Referral

Why the limits sit inside the professional ones

Professional mediators screen cases before accepting them, using structured instruments, to decide whether a dispute belongs in joint process at all. This playbook has no screening step, no trained neutral, and no ability to modify process for safety. Its limits therefore have to be drawn conservatively.

Violence and coercion

This deserves separating from everything else, because it is most likely to arise in exactly the settings this playbook targets — families, households, small organizations.

Where a relationship involves violence, coercion, or controlling behavior, joint collaborative process is not merely less effective. Advocates in the field hold that resolving disputes in such relationships through consensus-based process may be unsafe, unfair, and ineffective for the person at risk, and that these situations need the oversight and accountability a formal proceeding provides. Professional standards direct mediators to terminate when a participant cannot proceed safely and without fear or coercion.

A structured exercise can make this worse. It creates a record of what someone said. It can be used afterward. It can appear to give consent to an outcome that was not freely reached.

Do not use this playbook to work through a disagreement with someone you are afraid of.

The four responses

Green — continue. Substantive disagreement the group can hold. *“Your evidence doesn’t persuade me.” “I think you’re weighting that risk too heavily.” “I understand and still disagree.”*

Yellow — pause and reset. Heat without a safety concern: interruptions, sarcasm, profanity replacing explanation, talking points repeating without response, someone visibly overwhelmed. Call the stop word. Ask what you are actually disagreeing about. Take five minutes or switch to a simpler move. Still inside the playbook.

Orange — bring in a human facilitator. Participants are capable; self-facilitation is not working. Persistent power imbalance, repeated rule violations, entrenched interpersonal conflict, significant money or organizational consequences, accusations requiring adjudication, or lost trust in the neutrality of the process. Stop using GenAI as primary facilitator; continue only with an appropriate human mediator.

Red — stop. The difficulty is not really about the issue.

The stop word has a failure mode

A stop word gives control to whoever uses it, and that cuts both ways. The person with less power may not feel safe calling it. The person with more power can call it every time an uncomfortable line of questioning begins.

If one participant invokes the reset repeatedly and it is always the same subject, treat that as an Orange signal, not as good process discipline.

The Red list

Stop and refer where there are threats or fear of violence; coercion, intimidation, or controlling behavior; suspected abuse; serious emotional or behavioral crisis; self-harm or threats of harm; severe impairment; trauma reactions making continuation unsafe; any participant who cannot take part voluntarily; potentially criminal conduct; or legal rights and liability.

Again: you are not being asked to determine which of these is occurring. You are being asked to stop when you wonder.

Situation	Go to
Immediate danger	Emergency services
Suicidal thoughts, crisis	988 — call, text, or chat at 988lifeline.org
Domestic violence or coercive control	1-800-799-7233, National Domestic Violence Hotline
Relationship patterns, distress, trauma	Qualified counselor or therapist
Entrenched interpersonal dispute	Trained mediator
Workplace or organizational	HR, ombuds, or professional facilitator
Legal rights or liability	Attorney or formal dispute resolution

Resources are United States. Identify local equivalents before use elsewhere.

Standing notice

Reproduce with any exercise used outside its original context:

The Discussion Divide Playbook is a discussion and facilitation resource, not counseling, therapy, legal advice, crisis intervention, or a substitute for professional mediation. When safety, coercion, significant emotional distress, or serious power imbalance becomes part of the situation, stop the exercise and move to an appropriate human professional.

What this borrows, and from where

Mediation contributes neutral structure, reframing, working from interests rather than positions, voluntary participation, and participant-generated outcomes — a mediator asks questions and helps parties understand each other without taking sides or deciding, which

is the role this playbook assigns to GenAI. Positions-to-interests is from Fisher and Ury, *Getting to Yes*.

Counseling practice contributes active listening, reflection, boundaries, and establishing structure before attempting resolution.

Trauma-informed practice contributes the principle that psychological safety precedes finishing the exercise. SAMHSA's six guiding principles are safety; trustworthiness and transparency; peer support; collaboration and mutuality; empowerment, voice and choice; and cultural, historical, and gender issues. The fifth is why the stop word exists — it gives every participant real control, not a courtesy.

This playbook borrows process disciplines from these fields. It does not practice any of them.

Appendix B — Rules of Engagement

1. Agree on a stop word. Any participant, any time, no explanation. See Start Here and the failure mode in Appendix A.

2. Critique statements, not people. *Instead of:* "Anyone who believes that is stupid." *Try:* "I don't understand how that conclusion follows from the evidence."

3. Use words that carry information. Profanity and loaded labels communicate very little. *Instead of:* "That policy is garbage." *Try:* "I think the policy is unlikely to achieve its stated objective, because..." Strong feeling is legitimate. Convert it into information.

4. Describe before you judge. *Instead of:* "They're trying to suppress voters." *Try:* "This provision requires specific documentation, and I'm concerned people without it won't be able to register." *Instead of:* "They don't care about election security." *Try:* "Their proposal appears to weight access more heavily than front-end verification."

5. Do not assign motives. *Avoid:* "You only believe that because..." *Prefer:* "What leads you to that conclusion?" GenAI follows the same rule. It may offer possible interpretations. It may not state unverified motives as fact.

6. One claim at a time. Do not answer one disputed point with five new ones.

7. Separate facts from values. Evidence may settle a factual disagreement. It will not settle a disagreement about acceptable risk, fairness, liberty, cost, or priority.

8. Let people state their own position. The test: *would they say that is a fair description of what they believe?* If not, revise before debating. Applies especially to GenAI summaries.

9. No forced agreement. All successful outcomes: "We agree on the facts and weigh the tradeoffs differently." "We disagree about the evidence." "We agree on the goal and disagree about the method." "We don't have enough information yet."

10. Anyone may pass. No one must disclose personal experience, beliefs, family circumstances, political affiliation, or religion. "Pass" is a complete answer.

11. Stop means stop. See Appendix A.

12. Human judgment remains in control. Participants decide what evidence is credible, what values matter, what risks are acceptable, and what they believe.

Appendix C — Prompt Library

Fair summary. *We disagree about [topic]. Before evaluating our positions, ask each of us what we believe, what concerns us, what values we are trying to protect, and what evidence we consider credible. Summarize our positions in terms we both agree are fair.*

Shared goal. *Two people disagree about [topic]. Summarize the outcome each appears to be protecting. Identify any shared goal without pretending they agree on the method.*

Separate the layers. *Separate our statements into verifiable facts, assumptions, causal interpretations, forecasts, value judgments, emotional concerns, and policy preferences. Identify where we agree, where evidence can resolve the disagreement, and where different values or risk tolerances remain.*

Strongest other side. *Help each participant state the strongest legitimate version of the other's position. Flag caricatures, loaded wording, or claims the other person has not actually made.*

Falsification test. *Ask each participant what evidence, observation, or outcome would reasonably weaken their position. Do not pressure anyone to change their values.*

Evidence check. *Establish a shared fact base for [topic]. Weight each claim by the quality and directness of its evidence rather than how often it is repeated. Identify what remains uncertain and what we should independently verify.*

Reframe the exchange. *Review this exchange for talking points, deflections, false premises, unanswered questions, and claims carrying unequal evidentiary support. Reframe each as a constructive question we could examine together.*

Tradeoff table. *Build a neutral tradeoff table for these options, including doing nothing. Include benefits, costs, burden distribution, evidence strength, and uncertainties. Do not recommend a winner.*

Two-laptop mapping. *Using only the participants' own statements, produce two neutral summaries. Do not infer motives. Then map: agreement; factual disagreements evidence could resolve; value or risk-tolerance differences; unanswered questions.*

End-of-discussion update. *Summarize where we now agree, what factual questions remain, what evidence changed our understanding, what values still differ, and what would make the next discussion more productive.*

Appendix D — Evidence Notes

Supported. The Fair Summary rests on Eyal, Steffel and Epley (2018) — 25 experiments finding perspective-taking did not improve accuracy in understanding others and tended to

reduce accuracy while increasing confidence. One Real Story rests on Kalla and Broockman (2020), three field experiments where argument alone had no effect and added non-judgmental narrative exchange produced durable reductions in exclusionary attitudes lasting at least four months; and Broockman and Kalla (2016, *Science*). The Discussion Bridge structure is supported by Stanford's America in One Room.

Mixed. Test the Claims descends partly from Fernbach, Rogers, Fox and Sloman (2013) on the illusion of explanatory depth. Three preregistered replication failures exist, and the effect appears limited to consequentialist rather than protected-value issues.

Consistent. Find the Shared Goal, Sort the Statements, and Weigh the Tradeoffs align with deliberative practice but have not been tested as discrete techniques.

Framework. Coming Back is a Perpetual Innovation™ contribution.

On how this evidence exists. The deep-canvassing findings are trustworthy partly because of how the field handled a failure. A widely publicized 2014 study reporting very large persuasion effects was retracted after Broockman, Kalla and Aronow found the data had been fabricated. The same researchers then ran the experiment properly and found a real effect — far smaller than the fraudulent one, and durable. The people most invested in the finding were the ones who checked it, and they took a smaller true answer over a larger false one. That is the standard this playbook asks of its readers.

Sources

Eyal, Steffel & Epley (2018) —
<https://www.frontiersin.org/journals/communication/articles/10.3389/fcomm.2021.611187/full>

Broockman & Kalla (2016), *Science* —
<https://www.science.org/doi/10.1126/science.aad9713>

Kalla & Broockman (2020) —
https://papers.ssrn.com/sol3/papers.cfm?abstract_id=3532688

Fernbach, Rogers, Fox & Sloman (2013) —
<https://journals.sagepub.com/doi/abs/10.1177/0956797612464058>

Protected-values limitation —
<https://www.sciencedirect.com/science/article/abs/pii/S0010027722001342>

Braver Angels program evaluation — <https://braverangels.org/evaluation/>

Stanford America in One Room — <https://deliberation.stanford.edu/america-in-one-room>

SAMHSA, Six Guiding Principles — <https://www.samhsa.gov/resource/dbhis/infographic-6-guiding-principles-trauma-informed-approach>

NIST, Generative AI Profile — <https://www.nist.gov/publications/artificial-intelligence-risk-management-framework-generative-artificial-intelligence>

Appendix E — License

Original text, worksheets, prompts, and activity instructions are offered under Creative Commons Attribution-NonCommercial 4.0 International (CC BY-NC 4.0).

You may use, reproduce, teach, share, and adapt for noncommercial purposes with attribution. **Commercial use** — commercial training, consulting products, paid courseware, republication, or resale — requires separate permission from Strategic Business Planning Company.

The Discussion Divide Playbook, © 2026 Strategic Business Planning Company / PerpetualInnovation.org. Licensed under CC BY-NC 4.0.

Not legal advice; does not override rights in third-party content. Third-party programs listed in “Use What Already Exists” are the property of their respective organizations and are listed for reference only.

A better discussion does not have to end in agreement. It should end with a clearer understanding of what the disagreement actually is.