

Day 14 Presenter Notes

Shahryar Minhas

HOW TO USE THESE NOTES

Keep the teaching deck visible to students and use this file on your phone as the spoken route. Each section follows one slide. The slide content is summarized first, followed by the teaching note embedded in the deck. Do not run the models during class. The rendered walkthrough already contains the code and output, with the technical detail available whenever you want to slow down.

SLIDE 1: TITLE AND CONTINUITY

What to say

Frame this as the conclusion of the full course, not a new topic. We put the methods from Days 8 through 13 on common ground and decide which question each method is actually built to answer.

SLIDE 2: THE QUESTION THAT OUTLASTS THE COURSE

What students see

What is the substantive quantity you want, where does the important dependence live, and which assumptions are you willing to defend?

Choose a model from the question and data structure. Do not choose a question because a familiar model is available.

What to emphasize

Opening point: this is not a recap of package syntax. One rebel-cooperation network anchors the fitted model comparison so differences in estimand and interpretation stay visible. The frontier portion comes later, after we have finished reviewing the tools covered in the course.

SLIDE 3: DAYS 8 THROUGH 10 BUILT THE RELATIONAL TOOLKIT

What students see

Day 8: Network foundations

- Question: What exactly are the actors, ties, eligible pairs, and patterns in this network?
- Tool: netify for building, checking, reshaping, and describing network data.

Day 9: SRM and blockmodels

- Question: Are broad actor differences or recurring relational roles organizing the ties?
- Tools: lame and blockmodels for additive effects and cross-sectional roles, plus NetMix for actor-year role mixtures. The fitted Day 9 NetMix model did not impose a transition process between years.

Day 10: Latent distance, latent factor models, and AME

- Question: Which pairs remain unusually compatible after measured covariates and broad actor differences?
- Tools: latentnet for the distance model and lame for signed latent relational surfaces, covariate adjustment, uncertainty, and goodness-of-fit.

The sequence moved from defining the network, to locating broad actor structure, to representing pair-specific structure that the measured variables missed.

What to say

Begin with the practical foundation. On Day 8, we learned that a network model is only as coherent as the data object it receives. We decided who the actors were, what counted as a tie, which pairs could have ties, what a missing value meant, and whether the network was directed, weighted, layered, or observed repeatedly. netify was the common tool for building and checking those objects. This was not administrative setup. Every model later in the course inherited those choices.

Day 9 asked where broad relational differences live. The SRM gave actors additive effects, which let us ask who generally appeared more active or exposed after measured vari-

ables were included. Blockmodels asked a different question: whether a small number of recurring relational roles could summarize partner patterns. blockmodels handled the cross-sectional example. NetMix estimated a role mixture for each actor-year, but our specification did not impose a transition model connecting those mixtures over time. lame handled the SRM.

Day 10 moved from broad actor tendencies to pair-specific structure. The latent-distance model represented similarity through proximity, while the latent factor model allowed a signed pair-specific surface that could also represent complementary roles. AME combined measured predictors, additive actor effects, and that factor surface in one model. latentnet supplied the distance model, while lame supplied the AME fits, uncertainty when we used MCMC, and checks of whether the model reproduced important network patterns. The fast longitudinal examples used ALS point estimation instead, so those fits did not produce a posterior distribution.

The progression matters. Day 9 did not become obsolete on Day 10. Additive effects and discrete roles answer useful questions of their own. Day 10 added a continuous representation when those summaries were too coarse.

What to check for understanding

Ask students to distinguish these three statements:

- This actor has many ties across partners.
- This actor belongs to a recurring relational role.
- This particular pair relates more than the measured predictors and broad actor tendencies imply.

Those are the SRM, blockmodel, and latent factor questions in ordinary language.

SLIDE 4: DAYS 11 THROUGH 13 ADDED STRUCTURE, CHANGE, AND DESIGN

What students see

Day 11: ERGM and TERGM

- Question: Do named configurations such as reciprocity or closure help describe one network or a panel of networks?
- Tools: `ergm` and `btergm`, change statistics, simulation-based interpretation, and goodness-of-fit.

Day 12: SAOM

- Question: What actor-oriented sequence of small tie and behavior changes could connect the observed waves?
- Tools: RSiena with `netify`, method-of-moments estimation, micro-step interpretation, and simulation checks.

Day 13: Network causal inference

- Question: What comparison would identify the effect, and how can interference or shared actors break an ordinary design?
- Tools: exposure mappings, Horvitz-Thompson estimates, negative controls, and careful limits on causal claims.

These tools do not form a ladder from simple to advanced. They answer different questions about configurations, change between waves, and causal comparisons.

What to say

Day 11 made specific graph configurations part of the model. An ERGM asks whether a proposed tie is more or less compatible with the rest of the observed graph when that tie changes statistics such as reciprocity or shared partners. A TERGM carries configuration-based reasoning to repeated network observations. The practical tools were change statistics for interpretation and simulations for checking whether the fitted model could reproduce the network features we cared about.

Day 12 offered a different account of change. A SAOM imagines actors receiving opportunities to make small changes between the waves we observe. RSiena estimates parameters by simulating those unobserved paths and adjusting the parameters until selected simulated statistics match the observed statistics. We interpreted coefficients through specific candidate micro-steps, not as ordinary tie-level odds ratios.

Day 13 made us separate questions that are often blurred together. What is the causal comparison? How can one actor's treatment affect another actor's outcome? An exposure mapping and Horvitz-Thompson estimation handled the design-based examples. Negative controls and a GMM confounding bridge supplied a much more assumption-heavy observational strategy.

The broad lesson is that ERGM, TERGM, SAOM, and causal designs are not interchangeable versions of a dynamic network model. They encode different objects, processes, and identifying assumptions.

What to check for understanding

Ask what each of these would change:

- Adding a closure statistic to the conditional mean.
- Changing from a panel-level configuration model to an actor-oriented micro-step model.
- Keeping the fitted coefficients but replacing the ordinary covariance estimate with a dependence-aware one.

The answers are, respectively, the dependence representation, the temporal process, and the uncertainty calculation.

SLIDE 5: START BY NAMING THE TARGET

What students see

Possible targets include:

- An observed covariate association with calibrated uncertainty
- Actor sending and receiving propensities
- Discrete relational types
- A latent relational surface
- An explicit graph configuration such as closure
- Influence among actors

Is dependence something to account for, a process you want to describe, or the main thing you want to measure?

What to emphasize

First question: name one target and one nuisance feature. That keeps the payoff table from becoming a contest over

which model looks most sophisticated.

SLIDE 6: EVERY LENS MUST EARN ITS PLACE

What students see

For each model, finish four sentences:

1. **The substantive question is...**
2. **This model helps because...**
3. **It is distinct from the previous lens because...**
4. **After fitting it, I can say...**

If the fourth sentence still sounds like package documentation, the model has not yet produced an applied result.

What to emphasize

The walkthrough completes these sentences for GLM, SRM, blocks, AME, ERGM, TERGM, and SAOM. I will use the same four prompts as we move through the lenses, always naming the actors and recorded action before the method.

SLIDE 7: THE SAME NETWORK BECOMES DIFFERENT EVIDENCE

What students see

- **GLM.** Measured dyadic associations
- **SRM.** Broad actor differences
- **Blockmodel.** Recurring partner-list types
- **AME.** Continuous pair-specific compatibility

- **ERGM.** Named graph configurations

TERGM and SAOM require repeated network observations, so they belong in the course map but cannot be fitted to today's single snapshot.

What to emphasize

The pattern column says what each model is trying to explain. The same graph becomes different evidence because each cross-sectional procedure conditions on a different representation of network dependence. TERGM and SAOM are shown so students can see where the longitudinal tools fit, followed by the reason they are not part of the Gade application.

SLIDE 8: WHAT EACH PROCEDURE ACTUALLY FITS

What students see

- **GLM.** Maximize the dyad-level Bernoulli log-likelihood.
- **Capstone SRM and AME.** Sample a joint posterior distribution.
- **Blockmodel.** Maximize a variational likelihood lower bound.
- **Capstone ERGM.** Use MCMC maximum likelihood to maximize a graph likelihood.
- **TERGM with btergm.** Fit a conditional tie-transition model by maximum pseudolikelihood and use a temporal bootstrap for uncertainty.
- **SAOM.** Use simulated method of moments to match

observed network and behavior change statistics.

What to emphasize

Separate the statistical target from the search routine. A likelihood, posterior, variational bound, or moment condition defines what counts as a better fit. MCMC, variational EM, alternating updates, and simulation describe how the computer searches for or explores that target. The capstone ERGM uses MCMC maximum likelihood, or MCMLE, to approximate the graph-likelihood score with simulated networks. The `btergm` examples instead use maximum pseudo-likelihood, or MPLE, and a temporal bootstrap. TERGM and SAOM complete the course map even though this single snapshot cannot support either model. DCR is not listed as a fitted model because it has no new likelihood, posterior, moment condition, or coefficient target. It receives its own introduction after the GLM.

SLIDE 9: ONE DATASET KEEPS THE MODELS COMPARABLE

What students see

The rebel-cooperation network supplies:

- 30 armed organizations active in the Syrian civil war
- An undirected tie for at least one claimed tactical joint operation from July 2012 through June 2015
- Organization ideology and size
- Pair-level ideological distance, power difference, shared location, and shared sponsor

- Strong differences in how many cooperation partners organizations have

Every lens sees the same observed network, but each conditions on a different representation of its dependence.

What to emphasize

Rebuild the netify object and state the outcome, actor roster, direction, and missingness rule. The published paper models the square root of the number of claimed joint operations. This capstone deliberately changes the outcome to whether any claimed joint operation was recorded, so its coefficients are not numerical replications of the published estimates. Model comparison within the capstone is meaningful because all five fitted lenses use this same binary network.

SLIDE 10: LENS 0: START WITH THE DYADIC GLM

What students see

What it asks

Which measured pair characteristics are associated with recorded cooperation?

What it initially assumes

The usual GLM uncertainty calculation treats the 435 organization pairs as independent rows.

The GLM gives us familiar conditional associations. Before reading its intervals, we need to account for the fact that many rows contain the same organizations.

What to emphasize

Use one row per unordered organization pair and sum actor covariates across the two endpoints, matching the symmetric SRM, AME, and ERGM terms. Same location, ideological similarity, and combined group size are clearly associated with cooperation under the ordinary GLM calculation. Do not interpret the intervals yet. Point out that each organization appears in many rows, then use that overlap to motivate the next slide.

SLIDE 11: WHAT DYADIC CLUSTER-ROBUST STANDARD ERRORS DO

What students see

The regression rows overlap. ANF with ASIM and ANF with another organization both contain ANF. If something unmeasured makes ANF cooperate more widely, both regression errors can move together.

Dyadic cluster-robust standard errors, or DCRSEs, allow two errors to be related whenever the corresponding dyads share either organization.

DCR keeps every GLM coefficient fixed. It recalculates the standard errors, confidence intervals, and tests. Here, the power-difference interval widens across zero, while several other intervals narrow.

What to emphasize

Ordinary GLM uncertainty treats all 435 organization pairs as independent. DCR recognizes that pairs sharing ANF, ASIM, or another organization are not entirely separate pieces of information. It changes the uncertainty calculation, not the fitted associations and not the substantive model.

SLIDE 12: WHAT DCR DOES NOT FIX

What students see

- DCR does not add missing actor traits or network processes to the regression.
- It still treats dyads with no shared organization as independent.
- Its approximation can be unstable when the number of actors is small. This application has only 30 organizations.
- It does not turn an association into a causal effect.
- A corrected interval can widen or narrow because the estimated shared-actor relationships can have either sign.

DCR is an uncertainty correction for a regression we already chose. It is not a model of how the network formed.

What to emphasize

Do not present DCR as a substitute for SRM, AME, or ERGM. Those models change what the fitted relationship can represent. DCR asks a narrower question about uncertainty around an existing regression. Carlson, Incerti, and Aronow (2024) emphasize both the value of the shared-actor

correction and these limitations.

SLIDE 13: LENS 1: SRM LOCATES ORGANIZATION-LEVEL DEPENDENCE

What students see

$$\eta_{ij} = x_{ij}^T \beta + a_i + a_j$$

The SRM asks which organizations cooperate with more partners than the observed covariates predict, how much of the hub structure reflects organization-level heterogeneity, and which covariate results change after every organization receives an additive sociality effect.

What to emphasize

This undirected application has one overall cooperation effect per organization, not separate sender and receiver effects. ANF and ASIM have the largest fitted effects. The SRM reproduces the unequal partner counts without adding a separate friend-of-a-friend term. Compare the covariate results with DCR and locate whether disagreement comes from re-estimating the mean.

SLIDE 14: LENS 2: BLOCKMODELS ASK WHETHER ACTORS COME IN KINDS

What students see

$$\Pr(Y_{ij} = 1 \mid z_i, z_j) = \theta_{z_i z_j}$$

Can a small number of relational roles summarize which organizations cooperate?

The fitted two-block model places ANF and ASIM in a two-organization core and the other 28 organizations in one periphery. Its fitted probabilities are 0.96 for core-core, 0.74 for core-periphery, and 0.11 for periphery-periphery cooperation.

What to emphasize

The fitted two-block summary isolates ANF and ASIM as a small core and places the other 28 organizations in one sparse but heterogeneous periphery. Explain that this is a hub-versus-periphery summary, not two coalition communities. The hard partition contains only one possible core-core pair, so the 0.96 estimate is fragile and should not be read as a precise population rate. The full variational fit uses uncertain membership probabilities rather than only the hard labels.

SLIDE 15: LENS 3: AME ADDS A RELATIONAL SURFACE

What students see

$$\eta_{ij} = x_{ij}^T \beta + a_i + a_j + u_i^T \Lambda u_j$$

The SRM gives each organization one overall cooperation effect. AME adds a pair-specific score. A high score raises fitted cooperation for that pair, a low score lowers it, and the scores summarize recurring partner patterns left unex-

plained by the measured variables and overall organization effects.

Shared location remains clearly positive. The ideology-distance result becomes uncertain after the pair-specific surface is included.

What to emphasize

Explain the coefficient change in ordinary language. The measured ideology difference and the latent partner pattern compete to explain some of the same ties. AME therefore changes what the ideology coefficient means by adding pair-specific structure. The walkthrough reports coefficient effective sample sizes from the longer teaching chain. The displayed latent geometry is an exploratory point summary from one aligned posterior mean surface, not an uncertainty plot for latent positions. Keep coordinates separate from invariant fitted quantities, and use independently initialized chains for research.

SLIDE 16: LENS 4: ERGM MAKES CONFIGURATIONS EXPLICIT

What students see

$$\Pr(Y = y) \propto \exp\{\theta^\top g(y)\}$$

Does a shared-partner statistic retain a conditional association after actor heterogeneity is represented?

Use change statistics and simulation to interpret:

- Edges and covariates

- Actor sociality and degree-related terms
- Geometrically weighted closure

What to emphasize

State the question before the specifications: after accounting for organizations that cooperate unusually widely, is there still a clear shared-partner pattern? Hubs and closure can both create many shared partners, so the model must represent the hubs before the shared-partner term can answer that conditional question.

SLIDE 17: WHY HUBS CAN BE MISTAKEN FOR CLOSURE

What students see

ANF and ASIM cooperate with many organizations. That alone creates many pairs that share a cooperation partner, even if organizations did not choose partners through a friend-of-a-friend process. The first ERGM misses how unequal organizations' partner counts are and produces too much transitivity. Adding one curved degree term improves AIC but still misses both selected yardsticks. Giving each organization its own baseline activity covers the selected actor-heterogeneity and transitivity-dependence checks, and the remaining shared-partner interval includes zero.

What to emphasize

Explain the distinction directly. A hub pattern means a few organizations have many partners, which mechanically

creates shared partners. A closure pattern would mean ties remain especially common among pairs with shared partners after other included structure is represented. Do not compare the shared-partner coefficients numerically because each belongs to a different specification. A lower AIC for the middle model does not override failed simulation checks. The displayed goodness-of-fit comparison uses two scalar yardsticks, one for actor heterogeneity and one centered standardized third-order transitivity-dependence moment. It is not a full degree-distribution or shared-partner goodness-of-fit analysis. Read the final model on its own: after organization-specific activity is represented, these data do not show a clear additional shared-partner association. With 29 organization-specific sociality parameters and only 85 ties, treat this as a saturated teaching diagnostic, not proof that closure is absent.

SLIDE 18: WHAT DID EACH MODEL HELP US LEARN?

What students see

For each model, ask:

1. What was it trying to explain?
2. What feature did it add that the previous model lacked?
3. What did it reveal about these organizations and their cooperation?
4. What result or diagnostic would make us stop trusting it?

Do not ask which model won. Ask which model answers the research question and whether its assumptions and diagnos-

tics are credible.

What to emphasize

Use the table to translate each fitted model into an applied takeaway. The SRM helps us see broad differences among organizations, the blockmodel tests whether a few recurring partner types are useful, AME looks for pair-specific compatibility, and the ERGM tests named graph configurations. They are not contestants pursuing one identical target.

SLIDE 19: AFTER NAMING THE TARGET, PRESERVE THE DATA STRUCTURE

What students see

One snapshot

Cross-sectional

GLM/DCR, SRM, blocks, AME, ERGM

Panel waves

Longitudinal

TERGM, SAOM, dynamic latent models

Treatment spills over

Causal design

Day 13's exposure and comparison logic

What to emphasize

The estimand comes first. Next, preserve what was observed rather than forcing it into a familiar format. A snapshot, a panel, and a spillover design support different questions and require different assumptions. Multilayer data enter later as a frontier extension, not as a tool students were already expected to know.

SLIDE 20: CHOOSE A MODEL IN THREE STEPS

What students see

Start with what you want to learn, not with the package you want to use.

1. Say the question in ordinary language. Example: Which organizations cooperate unusually widely?
2. Choose the model feature that answers it. An SRM gives each organization its own overall cooperation tendency.
3. Name the check that could change your mind. If the SRM still misses important partner patterns, add a block-model, AME surface, or explicit ERGM configuration depending on the miss.

The next model should address a specific problem that the first model left behind.

What to emphasize

Walk through the example literally. If the question is broad organizational activity, start with the SRM. If simulated

networks still miss recurring partner types, try blocks. If they miss smooth pair-specific compatibility, try AME. If the theory names a configuration such as shared partners, try an ERGM and check it through simulation. Model expansion should follow a visible miss or a different research question.

SLIDE 21: FOUR STUDIES, FOUR LESSONS WE CAN REUSE

What students see

The voting-message study shows that an intervention can reach an untreated person through a household connection.

The school anti-conflict program distinguishes a student's own assignment, exposure through assigned peers, and the wider context of attending a program school.

The Ghana observer study shows that an intervention may move behavior to nearby untreated places.

The friends' GPA study shows that similar friends do not by themselves establish peer influence.

What to emphasize

These are not four isolated studies to memorize. They leave us with four habits: ask where an effect can travel, distinguish a unit's own assignment from exposure through assigned neighbors and the broader setting, look for behavior moving somewhere else, and ask why connected units already resemble one another.

SLIDE 22: CARRY SIX QUESTIONS FROM THE CAUSAL LECTURE

What students see

1. What changes? Name the treatment or exposure precisely.
2. Whose outcome may change, and when?
3. Compared with what? Name the units or periods supplying the comparison.
4. Why might treated and untreated cases already differ?
5. Can treatment reach a unit through its network neighbors?
6. Which result would still be useful if a causal claim is not credible?

Do this before choosing a network model. A sophisticated dependence model cannot repair a comparison that never identified the effect.

What to emphasize

Most students will use observational data. The practical lesson is to make the comparison and its weaknesses explicit before fitting anything. The network may define the treatment exposure, the process that selected units into treatment, dependence in the outcomes, or the outcome itself. Those are different jobs.

SLIDE 23: TURN THE QUESTIONS INTO A WORKABLE ANALYSIS

What students see

Example: Do sanctions reduce later trade?

- Define sanctions before the trade outcome.
- Compare state-pair periods with credible overlap.
- Measure pre-sanction trade, security relations, regime ties, and economic conditions.
- Ask whether allies' sanctions also affect the focal pair. If so, define that network exposure.
- Use pre-treatment outcomes or negative controls to probe hidden selection.
- Model repeated states and dyads, but keep dependence separate from identification.

If the conditions producing sanctions may also reduce trade, report an association and name that remaining threat.

What to emphasize

Translate this to other applications. In election monitoring, ask why observers went to some locations. In policy diffusion, define which neighbors count and how their adoption becomes exposure. In conflict spillovers, name the geographic, alliance, or organizational channel. Latent positions may help proxy omitted traits when those traits leave a recoverable footprint in the network, but they do not manufacture a credible comparison.

**SLIDE 24: HEADACHES ARE A CHECK,
NOT THE OUTCOME**

What students see

Egami and Tchetgen Tchetgen ask whether students do better academically when their friends have higher GPAs.

The problem is that students choose their friends. Shared motivation, family resources, stress, or school conditions could make friends' GPAs look influential even when the relationship comes from shared background.

Friends' headaches should not directly change the student's later GPA. They may, however, reflect that hidden background.

What to emphasize

Keep the roles straight. The outcome is the student's later GPA. Friends' GPA is the exposure. Friends' headaches are a diagnostic clue, not another outcome the researchers care about.

SLIDE 25: WHAT THE TWO NEGATIVE CONTROLS TELL US

What students see

Friends' headaches give a clue about hidden features of the friendship environment but should not directly change the student's later GPA. The student's earlier GPA gives a clue about background influences on academic performance but cannot be changed backward in time by friends' GPA.

If friends' headaches still predict later GPA, or friends' GPA appears to predict a GPA recorded earlier, then unmeasured

network confounding may remain or the proposed negative control may violate its exclusion assumptions.

What to emphasize

First use negative controls as a warning check. Under stronger assumptions, Egami and Tchetgen Tchetgen use both clues to help adjust the peer-GPA estimate. This is not the same as simply adding headaches and earlier GPA to an ordinary regression.

SLIDE 26: THE FRONTIER: PREDICTING COOPERATION FROM NETWORK POSITION

What students see

Return to the 30 armed organizations in the Syrian civil war and their 85 recorded cooperation ties.

- Al-Nusra Front, or ANF, cooperates with 20 organizations.
- Ahrar al-Sham Islamic Movement, or ASIM, cooperates with 25.
- ANF and ASIM share 18 cooperation partners.

We repeatedly hide cooperation ties before building any network features and ask which model best recovers them. Degree and shared-partner features produce a mean AUC of 0.82. Walk-based coordinates produce 0.60. Combining both produces 0.80. An AUC of 0.50 is chance and 1.00 is perfect ranking.

The simple network summaries outperform the walk-based coordinates in this small, hub-dominated network. Node2vec turns an actor's network position into a short list of predictors, but it does not automatically improve prediction or explain why cooperation occurred.

What to say

This is where the frontier portion begins. Node2vec was not one of the models covered earlier in the course. I am introducing it now as an example of how machine learning researchers use a network to build predictors.

Keep the application concrete. We have 30 armed organizations and 85 recorded cooperation ties. Al-Nusra Front, which I will call ANF, has 20 cooperation partners. Ahrar al-Sham Islamic Movement, which I will call ASIM, has 25. They share 18 partners, so their places in the network overlap considerably even before we decide how to summarize that overlap.

The prediction exercise is fitted. In each of 30 splits, we hide 17 recorded cooperation ties and 17 non-ties before constructing any features. We compare a simple baseline based on degree and shared partners with walk-based coordinates and with a model that combines both. The baseline has mean AUC 0.82, compared with 0.60 for the coordinates and 0.80 for the combined model. Explain AUC as a ranking measure: if we show the model one hidden tie and one hidden non-tie, 0.82 means it ranks the tie higher about 82 percent of the time. The teaching implementation uses node2vec's biased walks and then compresses a PPMI neighborhood table with SVD. It demonstrates the

walk-embedding idea, but it is not the original skip-gram implementation.

The important shift is the goal. Earlier models tried to describe or explain a declared network outcome. Here the immediate goal is to turn each organization's network position into useful predictors for a separate prediction task. The negative result is useful: a sophisticated representation has not earned its complexity when a simple baseline predicts better.

SLIDE 27: WHAT NODE2VEC DOES WITH THE NETWORK

What students see

1. Reserve ties and non-ties for testing, then remove the test ties from the graph.
2. Start at one organization and follow a few observed ties.
3. Repeat these short network trips from every organization.
4. Give organizations similar numerical profiles when the trips repeatedly reach similar neighborhoods.
5. Use those profiles in a prediction model.

The trips are created by the computer. They are not organizations moving through the conflict.

The result is a few numbers for each organization. Those numbers are useful features, not automatically named traits such as ideology or influence.

What to say

Walk through the five steps slowly. First, reserve a test set containing ties and non-ties, then remove the test ties before creating any features. If those ties stay in the network while we build the features, the algorithm has already seen part of the answer and the evaluation is contaminated by leakage.

Second, the computer starts at one organization and follows a few recorded cooperation ties. Third, it repeats these short network trips many times and starts them from every organization. A possible trip could move from one organization to ASIM, then to ANF, and then to another ANF partner. This is only a computer-generated route through recorded ties. It is not a sequence of events and does not mean that anything traveled through the conflict network.

Fourth, the algorithm places every organization at a point in a small learned space. The short row of numbers is the organization's coordinates, much like factor scores. Organizations encountered in similar network neighborhoods tend to receive nearby coordinates. Fifth, we combine the coordinates of two organizations in a prediction model and ask whether they help predict a hidden cooperation tie. In this fitted exercise, those coordinates do not outperform degree and shared-partner features.

The output is a short row of coordinates for each organization. We can use distances, dot products, or pairwise combinations of those coordinates as predictors. The algorithm does not tell us that the first axis is ideology or that the second is influence. Like latent-factor coordinates, the axes can rotate or change signs without changing the relationships they encode. We therefore interpret relative positions

and held-out performance, not isolated coordinate labels.

No knowledge of text analysis is needed. The historical connection is only that node2vec borrowed a machine-learning trick for learning from repeated contexts. In this application, the contexts come from computer-generated trips through a network.

SLIDE 28: THREE METHODS, THREE DEFINITIONS OF SIMILAR

What students see

- DeepWalk uses ordinary short trips. Actors repeatedly appearing in similar trip contexts receive similar coordinates.
- LINE learns directly from observed ties and from similarity between neighbor lists.
- Node2vec lets the trips stay near the starting actor or explore farther outward. Its settings determine whether similarity emphasizes nearby partners or broader structural roles.

The settings matter because similar network position can mean sharing nearby partners or playing a similar role in different parts of the graph.

What to say

DeepWalk introduced the straightforward walk-based approach. LINE was designed to learn from ties and neighbor patterns at very large scale without depending on long walks. Node2vec made the walks adjustable. A local setting

may place ANF and ASIM close because they share many nearby partners. A more exploratory setting can instead emphasize organizations that play similar hub or brokerage roles in different portions of a network. There is no universally correct setting because the prediction task determines which similarity is useful.

SLIDE 29: NETMF AND GRAPHSAGE SOLVE TWO PRACTICAL PROBLEMS

What students see

NetMF exposes the matrix. Qiu et al. show that popular embedding methods implicitly compress a large table of transformed node-context co-occurrences. NetMF writes down and factorizes that table directly. The connection to SVD and Day 10 is genuine, but the table is not the raw adjacency matrix.

GraphSAGE learns a reusable rule. It combines an actor's traits with summaries of its observed neighbors. The rule can construct coordinates for a new actor when those inputs are available.

NetMF clarifies what is being compressed. GraphSAGE changes the goal from storing actor coordinates to learning how to construct them.

What to say

Qiu et al. provide the clean bridge to the factor-model intuition. The matrix contains transformed information about which actors occur in which sampled network contexts. A

low-rank factorization gives each actor a short coordinate row. This resembles the SVD logic from Day 10, but it reconstructs a different object and serves a prediction-first goal.

GraphSAGE solves a different problem. Standard node2vec stores coordinates for actors seen during training. GraphSAGE learns an aggregation rule that can be applied later. A new actor still needs observed attributes and an observed neighborhood. The method cannot create useful information for an isolated, completely unobserved actor.

SLIDE 30: QUESTION 1: CAN WE PREDICT A HIDDEN TIE?

What students see

For the Syrian cooperation example:

1. Reserve ties and non-ties before learning coordinates.
2. Remove the reserved ties and learn node2vec on the reduced graph.
3. Build pair-level predictors from the two organizations' coordinates.
4. Fit on visible pairs and predict only the hidden pairs.
5. Compare with degree and shared-partner baselines. In a full application, also add measured pair characteristics.

A useful result would show improved prediction of held-back cooperation ties beyond those baselines. The compact teaching fit compares network-derived features with one another and does not include ideology, power, location, or sponsorship. A research application should add those measured

pair characteristics before claiming that the coordinates add information beyond observed predictors. Prediction would still not explain why the ties formed.

What to say

The ordering prevents information leakage. If a held-back tie stays in the graph used to learn the coordinates, the feature construction has already seen part of the answer. Because non-ties outnumber ties, raw accuracy can look good when a model predicts almost everything as a non-tie. Report ranking measures such as area under the ROC curve or precision among the highest-scored candidate pairs, and explain how non-ties were sampled.

SLIDE 31: QUESTION 2: CAN NETWORK POSITION PREDICT AN ACTOR LABEL?

What students see

Brown et al. (2021) use linking patterns to help predict ideological labels for Twitter users and internet domains.

- The outcome belongs to an actor, not a pair.
- Keep test actors in the graph so the transductive embedding can represent them, but hide their labels from the classifier.
- Compare network features with label-frequency and useful non-network baselines.
- Check performance across samples, platforms, or periods.

The warranted claim is that link neighborhoods contain pre-

dictive information about held-out labels. A coordinate does not automatically become an ideology scale.

What to say

Make the unit change explicit. Hidden-tie prediction holds out dyads before building network features. Transductive actor-label prediction keeps test actors in the graph so node2vec can embed them, but hides their labels while the classifier is trained. Removing complete actors would test the separate inductive GraphSAGE task on the next slide. Prevent leakage from labels or links recorded after the outcome period. Good classification does not show that network position caused ideology.

SLIDE 32: QUESTION 3: CAN WE BUILD FEATURES FOR A NEW ACTOR?

What students see

Ordinary node2vec has no stored coordinate row for an actor absent during training. GraphSAGE learns a rule that starts with the new actor's attributes, summarizes its neighbors, combines those inputs, and constructs coordinates for the prediction task.

Test on complete actors or later periods excluded from training. The method still needs observed attributes and enough neighborhood information for each new actor.

What to say

This is an inductive task because the actor was absent when the rule was learned. Holding out random dyads among actors already in training does not test this ability. Hold out complete actors or train on earlier periods and evaluate actors appearing later. Use only attributes and neighbors that would have been available at prediction time.

SLIDE 33: NODE2VEC AND DAY 10 ARE COUSINS

What students see

Node2vec and related graph machine learning build useful features for prediction by compressing network neighborhoods. Latent factor models and AME explain a declared dyadic outcome by compressing unexplained patterns in that outcome.

Both approaches reduce complicated relational patterns to a few numbers per actor. Node2vec is checked with genuinely hidden cases. AME is checked with uncertainty, prediction, and network goodness-of-fit.

If node2vec improves hidden-tie prediction, say exactly that. Do not rename an axis “ideology,” “influence,” or “alliance” without outside evidence.

What to say

Connect this directly to the movie example from Day 10. There, we used a few latent dimensions to summarize a much

larger table of viewer and movie relationships. The basic compression intuition is similar here: represent complicated relational information with a small number of values.

The research targets are different. Node2vec compresses neighborhoods into features for a prediction task. We ask whether those features work on ties or labels that were genuinely hidden while the features were built. AME models the declared dyadic outcome itself. It can include observed covariates and broad actor effects, and it gives us model-based uncertainty and network goodness-of-fit checks.

Do not turn the node2vec coordinates into an unsupported substantive story. If a fitted node2vec model improves prediction of hidden cooperation ties, the warranted statement is that the network features improved that held-out prediction. It would not establish why the organizations cooperated, identify influence, or make a coordinate an ideology scale.

Return to the fitted Syrian cooperation result. In this small network, the walk coordinates did not beat degree and shared-partner summaries on hidden ties. That is a useful finding because the machine-learning representation has to earn its complexity on the prediction task rather than receive credit for sounding more advanced.

SLIDE 34: SIR ASKS HOW CONFLICT PATTERNS CARRY FORWARD

What students see

Minhas and Hoff (2026) study monthly material-conflict events among countries. These include physical attacks, destruction of property, and other coercive actions recorded by ICEWS.

The model separates two questions:

- Does a state pair's own earlier conflict help predict its later conflict?
- Does conflict elsewhere in the international system help predict who a state targets next and who is targeted next?

Influence here means a lagged predictive association. It does not by itself show imitation, coordination, or a causal effect.

What to say

Begin with the outcome and time scale. Each monthly network records how many material-conflict events one country directed toward another. The model asks where the next month's conflict pattern comes from.

The direct part stays within one pair. If state i attacked state j last month, does that pair record more conflict this month? The influence part reaches beyond that pair. It asks whether the targets of i 's allies, or the sources that attack states near j , help predict the current i -to- j count.

This use of influence is narrower than a causal story. It means an earlier pattern elsewhere in the network carries predictive information for a later dyadic outcome after the other modeled terms are included. We would need a sepa-

rate design and stronger assumptions to say that one country caused another to act.

SLIDE 35: DIRECT EFFECTS AND INFLUENCE PATHS ARE DIFFERENT

What students see

The direct part includes the pair's earlier conflict, conflict in the reverse direction, geographic distance, joint democracy, alliance, trade, and verbal cooperation.

The influence part asks two different network questions:

- Sender side: Does state i target countries that its allies or verbal-cooperation partners targeted previously?
- Receiver side: Are geographically nearby states targeted by similar sets of countries?

What to say

Use one imagined dyad, i directing conflict toward j . The direct predictors describe i and j themselves, their relationship, and their own recent history. Earlier i -to- j conflict and earlier j -to- i conflict belong here, as do distance, alliance, trade, democracy, and verbal cooperation.

The sender-side influence path brings in a third state. If one of i 's allies targeted j last month, does that make current i -to- j conflict more predictable? The receiver-side path brings in similarities among targets. If countries near j were targeted by a particular set of senders, does j tend to receive conflict from those same senders?

The distinction matters because a large direct-lag coefficient and a large influence-path weight are not the same result. One says the dyad's own conflict persists. The other says a wider relational channel organizes where conflict appears next.

SLIDE 36: WHAT THE PUBLISHED SIR APPLICATION FINDS

What students see

- Alliances organize sender-side influence: states tend to direct more conflict toward countries that their allies fought in the previous month.
- Verbal cooperation also organizes sender-side influence: states tend to target some of the same countries as their cooperative diplomatic partners.
- Geography organizes receiver-side influence: nearby states tend to receive conflict from similar sets of senders.

These are conditional temporal patterns. The model does not establish that allies coordinated, that one state copied another, or that proximity caused conflict.

What to say

Translate each result into a concrete pattern. First, the targets of a state's allies carry information about that state's next targets. Second, the targets of states with which it cooperates verbally also carry information. Third, countries that are geographically near one another tend to be attacked

by overlapping sets of senders.

The model is useful because it does more than say conflict is temporally dependent. It tells us which observed relational channels best organize that dependence. Keep the claim on the scale the model supports: these are lagged conditional associations in monthly ICEWS counts. The fit does not show whether allies deliberately coordinated or whether one state copied another.

SLIDE 37: THE SAME MODELING MOVE CAN STRUCTURE DIFFERENT OBJECTS

What students see

NetMix, the covariate-defined latent space in Austin et al. (2013), and SIR all connect observed covariates or relations to an unobserved relational object. They do not estimate the same object.

What to emphasize

In NetMix, nodal covariates help predict an actor-year's mixture of discrete roles. In Austin et al., nodal covariates define an actor's expected continuous latent position, while a random residual allows departure from that expectation. In SIR, observed relational matrices define possible directed influence channels over time, and estimated weights describe which channels carry the lagged association. The shared lesson is that observed information can organize a hidden relational structure. Do not interpret the resulting roles,

positions, or influence channels as interchangeable.

SLIDE 38: HOW SIR IS ESTIMATED

What students see

For the monthly conflict counts:

$$\eta_{ijt} = z_{ijt}^T \theta + (A X_{t-1} B^T)_{ij}, \quad A = \sum_r \alpha_r W_r, \quad B = \sum_r \beta_r W_r$$

- $z_{ijt}^T \theta$ contains direct dyadic predictors.
- X_{t-1} is the previous month's conflict network.
- A describes sender-side channels and B describes receiver-side channels.

The Poisson estimator holds one side fixed, updates the other side with a familiar GLM step, and repeats until the fitted likelihood and coefficients stabilize.

What to emphasize

The Poisson likelihood chooses coefficients that make the observed monthly counts most plausible under the model. Holding the receiver-side coefficients fixed turns the sender-side update into a familiar GLM job. Holding the sender side fixed does the same for the receiver side. The algorithm alternates until the likelihood and coefficients stabilize.

The separate alpha and beta scales are not individually identified, so the implementation normalizes one coefficient for reporting. The resulting channel product and fitted means

remain interpretable. The default sandwich covariance allows score contributions involving the same endpoint actor to move together. It does not cluster by month and does not change the fitted coefficients. The walkthrough's switch-off calculation is a fitted-prediction ablation, not a causal share or variance decomposition.

SLIDE 39: MULTILAYER NETWORKS KEEP SEVERAL RELATIONS VISIBLE

What students see

$$Y_{ij\ell}$$

records relation type ℓ from source i to target j among the same actors.

With repeated years, the multilayer object becomes Y_{ijlt} .

Examples include military threats, sanctions, trade, and alliances among states; four ICEWS cooperation and conflict categories; and advice, friendship, and collaboration among people.

What to emphasize

Begin with the applied reason for preserving the layer index. A state can cooperate with another state verbally, cooperate materially, criticize it, and engage in material conflict during the same period. Calling all four relationships “interaction” would erase the outcome distinction before any model is estimated.

The extra index ℓ tells us which relation produced a tie. If

time is added, t records when it occurred. A multilayer analysis can then ask whether the same states and state pairs remain prominent across relations, whether some patterns are specific to conflict or cooperation, and which residual profiles the layers share.

SLIDE 40: THEORY CAN MAKE THE LAYERS INTERDEPENDENT

What students see

Repeated multilayer data can represent several relationships changing together. In international relations, conflict can reduce trade while trade dependence may discourage later conflict. In American politics, cosponsorship can build voting coalitions while voting agreement may make later cosponsorship easier. In education, friendship creates opportunities for study help while repeated help can strengthen friendship.

What to emphasize

Emphasize the two-way possibility in each example. The point is not merely that two relations exist. The theory says one relation may shape the other and the reverse may also occur. A multilayer design preserves both outcomes so we can study their sequence and shared structure. Repeated observations help establish sequence, but they do not by themselves establish causal direction.

SLIDE 41: WHAT NETIFY CAN DO WITH THE LAYERS

What students see

netify can combine aligned networks, preserve layer labels, check and summarize the object, extract a layer or matrix, compare layers, and plot them through one common data structure.

What to emphasize

The practical payoff is safe data handling. A common actor order and explicit layer labels reduce silent mismatches. **compare_networks()** supplies descriptive correlations and overlap measures, while extracted layers can enter layer-specific models and the aligned collection can become the array needed by a joint model. The software does not decide whether a joint model is theoretically warranted.

SLIDE 42: ONE NETWORK CAN CONTAIN FOUR DIFFERENT STORIES

What students see

Four ICEWS panels show verbal cooperation, material cooperation, verbal conflict, and material conflict among the same 12 high-activity states in 2010.

The state positions are fixed across panels. Each panel displays its 12 highest-volume directed state pairs, and arrow thickness is scaled within that layer.

What to say

Read across panels while keeping the positions fixed. This makes changing ties, rather than a changing layout, carry the comparison.

The United States is prominent across the four displayed layers, but its partners and the surrounding pattern change. Some state pairs appear in several panels, while others are visible only for one kind of interaction. Aggregating the relations would hide those differences.

Be precise about what the graphic can show. Each panel selects its own 12 highest-volume pairs and scales arrow thickness within that layer. Compare which partners appear and how the pattern changes. Do not use this display to compare raw density or event volume across panels.

This is descriptive evidence. It does not establish that activity in one layer causes activity in another or that the four panels form a temporal sequence.

SLIDE 43: DO NOT FLATTEN AWAY THE QUESTION

What students see

Combining verbal cooperation, material cooperation, verbal conflict, and material conflict into “any interaction” changes the outcome. Four separate fits preserve the labels but cannot estimate shared source profiles, shared target profiles, layer loadings, or shared cross-layer mean structure.

What to emphasize

Flattening may be defensible if “any recorded interaction” is genuinely the outcome. It is not defensible if the question asks whether cooperation and conflict share a dyadic profile. Separate fits answer four useful within-layer questions but cannot estimate a relational mean structure shared across layers because they contain no joint layer parameter.

SLIDE 44: FOLLOW ONE SHARED COMPONENT ACROSS THE LAYERS

What students see

For component r , one source-target-layer contribution is $u_{ir}v_{jr}w_{\ell r}$. The source and target coordinates identify dyads matching the profile, and the layer loading says how strongly each relation expresses that same profile.

What to emphasize

Build the intuition one multiplication at a time. A high source coordinate and high target coordinate make a dyad strongly match a component. The layer weight can make that component strong, weak, reversed, or absent in a particular relation. In the teaching fit, one component emphasizes dyads involving the United States and another emphasizes a Middle East-centered set. Both recur in cooperation and conflict, so the fit does not support naming one component “cooperation” and the other “conflict.”

SLIDE 45: A JOINT MULTILINEAR MODEL SHARES ACTOR PROFILES

What students see

$$\eta_{ij\ell} = x_{ij\ell}^\top \beta_\ell + a_{i\ell} + b_{j\ell} + \sum_{r=1}^R u_{ir} v_{jr} w_{\ell r}.$$

The full model can include relation-specific measured predictors. The teaching fit omits that block and models separately standardized $\log(1 + \text{count})$ layers.

What to emphasize

Walk left to right. Beta allows a measured predictor to have a different association in each relation. The additive effects allow a state's overall activity and exposure to differ by layer. U and V are source and target profiles shared across relations. W says how strongly each relation expresses each shared component. W is the piece that does not exist in four separate fits. Then state the exact teaching simplification: the covariate block is omitted, and each event-count layer is transformed with $\log(1 + \text{count})$ and standardized separately before fitting.

SLIDE 46: WHAT THE TENSOR ESTIMATOR OPTIMIZES

What students see

For the transformed Gaussian example, the estimator minimizes

$$\sum_{i \neq j, \ell} \left(Y_{ij\ell}^* - \eta_{ij\ell} \right)^2.$$

It alternates updates for U, V, W, and the additive effects. This example compares eight starting values.

What to emphasize

Name the objective before the acronym. Minimizing squared residuals is maximizing a Gaussian working likelihood. Holding two factor blocks fixed makes the third update a familiar least-squares problem, but the full objective is not jointly convex. That is why starts matter. The implementation uses a small ridge term and short additive-effect updates, so say that each block is updated to reduce the common objective rather than claiming an exact unpenalized solution at every step. This classroom fit compares eight starts. A research fit should use more starts and select rank with held-out data. Binary and count tensor models change the likelihood and use weighted, variational, or Bayesian updates.

The individual coordinates can change sign, scale, and order without changing the fitted tensor. Interpret the summed fitted surface, normalized layer loadings, held-out predictions, and replicated network features.

SLIDE 47: OVERLAP NEEDS A DENSITY BENCHMARK

What students see

Verbal conflict and material conflict have Jaccard overlap of 0.385 at the one-event threshold and 0.402 at the five-event threshold. Their overlap lifts are 4.91 and 16.39.

Overlap lift compares the observed joint-tie rate with the joint-tie rate expected under dyad-level independence while preserving both layer densities.

What to emphasize

Jaccard answers how much two binary layers share. Lift asks whether that overlap exceeds what their densities alone would generate. The result survives the stricter threshold, which makes it less dependent on one arbitrary binarization rule. It still does not distinguish escalation from common news attention, coding practices, or persistent high-activity dyads.

SLIDE 48: WHAT SHOULD WE DO AFTER FINDING LAYER OVERLAP?

What students see

Check thresholds and coding, use a stronger null, match the next model to the claim, and inspect what remains unexplained.

What to emphasize

Overlap is a beginning, not an endpoint. First check whether the pattern survives defensible measurement

choices. Next ask whether it remains after preserving actor activity, time patterns, or other obvious sources of co-occurrence. Use a joint factor model when the goal is shared profiles, and repeated layers with lags when the goal is cross-layer prediction. A causal claim still needs a design beyond the overlap statistic.

SLIDE 49: DO THE SAME STATE PAIRS STAND OUT ACROSS RELATIONS?

What students see

After removing each state's broad source and target activity within each layer, the rank-2 surface accounts for about 23 percent of the remaining transformed variation.

Pattern 1 contains particular United States dyads that remain unusually active after the United States' overall activity is removed. Pattern 2 contains pairs involving Iran, Israel, Lebanon, Egypt, and Iraq. Both profiles appear in cooperation and conflict rather than producing a clean cooperation-versus-conflict axis.

What to emphasize

This is the useful result: the same residual pair profiles recur across several kinds of recorded interaction. The United States pattern is not simply high United States activity because the layer-specific source and target effects were removed first. Do not force a clean cooperation-versus-conflict story that the loadings do not show. The 23 percent is a reduction in squared residual variation after additive effects,

not an R-squared for raw event counts. Rank 2 is a compact in-sample point fit, not a selected truth. A research analysis needs covariates, rank sensitivity, held-out prediction, parameter uncertainty, and layer-specific goodness-of-fit.

SLIDE 50: HOW SHOULD WE CHECK THE JOINT FIT?

What students see

Compare ranks on held-out cells, repeat ALS from several starts, add defensible covariates, inspect residual network patterns, and use an outcome model and uncertainty procedure that match the data.

What to emphasize

The current rank-2 result is a useful teaching summary, not a selected truth. Held-out performance asks whether the extra rank generalizes. Eight starts are enough to demonstrate why the nonconvex objective needs repeated starts, but a research analysis should use more. Covariates show whether the shared profiles survive measured explanations. Residual checks ask what structure the fit missed. A binary or count outcome requires the corresponding likelihood, and research claims about uncertainty require a bootstrap or Bayesian fit.

SLIDE 51: SCALING SHORTCUTS CHANGE THE STATISTICAL JOB

What students see

ALS uses conditional least-squares updates but does not supply a posterior or protection from local optima. The capstone ERGM uses MCMC maximum likelihood, which targets the graph likelihood but can be computationally expensive and requires adequate simulation mixing. **btergm** uses MPLE plus a temporal bootstrap, which does not supply exact graph-likelihood estimation. Variational blockmodels optimize a lower bound. Bayesian AME uses MCMC posterior sampling. Graph embeddings use sampled walks and negative examples rather than a likelihood for the original tie outcome.

What to emphasize

Ask three questions of every shortcut: what objective does it optimize, what information does it sample or approximate, and does its uncertainty match the intended claim? Connect ALS directly to **lame**. Distinguish its blockwise updates within one squared-error objective from **btergm** pseudolikelihood, which replaces the graph likelihood with a product of conditional tie contributions. The capstone ERGM instead uses simulated networks for MCMLE. Speed is not a diagnostic of statistical adequacy.

SLIDE 52: WRITE THE MODEL CHOICE AS AN ARGUMENT

What students see

Use this paragraph structure:

1. **Question:** Name the estimand.
2. **Data:** State the network, time, and eligibility structure.
3. **Primary model:** Explain where it represents dependence.
4. **Runner-up:** Name the strongest alternative representation.
5. **Flip condition:** Give a specific diagnostic that would change the choice.
6. **Claim:** State what the fitted result would and would not establish.

What to emphasize

Tell students this is a disciplined way to justify a model, not a form to complete mechanically. The runner-up prevents them from presenting the chosen model as inevitable. The flip condition names evidence that would make them reconsider. The next three slides show the structure in published work.

SLIDE 53: PUBLISHED EXAMPLE: WITH WHOM DO SYRIAN REBELS COOPERATE?

What students see

Gade et al. (2019) ask why armed organizations chose some partners rather than others during the Syrian civil war. They use claims of tactical joint operations from July 2012 through June 2015, combined with measures of ideology, power, state sponsorship, and shared location. The compet-

ing explanations are that rebels choose ideologically similar partners, partners of similar strength, or partners backed by the same state.

What to emphasize

Connect the paper to the application students have already seen. The published AME models the square root of the number of claimed joint operations, while the classroom comparison uses a binary indicator for whether any claimed operation was recorded. The outcome difference is why the published and classroom ideology results need not match. Shared location matters because organizations cannot conduct joint operations together if they do not operate in the same area.

SLIDE 54: WORK THROUGH THEIR MODEL-CHOICE ARGUMENT

What students see

The question is which group differences are associated with more joint operations. The primary model is AME regression, which represents broad organization activity and remaining higher-order dependence. A strong model runner-up would be an ERGM with comparable covariates, organization heterogeneity, and theoretically named configurations, although the paper does not fit that model. The paper's activity-constrained simulations are a robustness check based on observed joint-operation counts, not another fitted coefficient model. Reconsider the ideology conclusion

if it disappears when shared location enters or if activity-constrained simulations reproduce it without ideology. The supported claim is that ideologically closer groups cooperated more often, not an identified explanation of why they chose one another.

What to emphasize

Walk down the framework using language students can reuse. State the outcome, data structure, dependence problem, primary model, strongest alternative model, published robustness check, evidence that would weaken the conclusion, and bounded claim. Keep the two comparisons distinct. An ERGM with comparable covariates and actor heterogeneity would be the true model runner-up because it represents dependence differently. The paper's simulations are an activity-constrained robustness check, not another estimator and not an AME coefficient comparison. They ask whether the broad ideology pattern remains when the observed operation activity is preserved.

SLIDE 55: WHAT THE PUBLISHED COMPARISON CHANGED

What students see

Ideological proximity had the clearest and most consistent association with cooperation across the published AME specifications. Shared location also mattered. Evidence for power similarity was weaker and less consistent. Shared state sponsorship did not receive clear support. Activity-

constrained simulations also showed more ideological similarity than activity levels alone would generate.

What to emphasize

The substantive conclusion is easy to state: organizations with more similar ideologies recorded more tactical joint operations in this fragmented conflict. Power and sponsorship did less to distinguish partnerships in these analyses. The AME regression and the activity-constrained simulation reach compatible descriptive conclusions through different routes. The published result and the capstone binary AME interval are not contradictory because their outcomes and conditioning sets differ. Neither analysis shows that changing an organization's ideology would change its partners.

SLIDE 56: THE MAP, FILLED IN

What students see

Build and describe

netify, summaries, centrality, communities

Actors and latent structure

SRM, blocks, latent distance, latent factors, AME

Graph configurations

ERGM and TERGM

Actor-oriented change

SAOM

Identification and uncertainty

DCR, exposure mappings, Horvitz-Thompson estimators, negative controls

Frontier tools

Node2vec, tensors, SIR, multilayer models

No cell replaces the others. Each makes a different part of relational structure explicit.

What to emphasize

Use this as the second course recap. The first row reminds students that the tools begin before estimation: netify and descriptive analysis define what the network is. The next rows separate latent structure, graph configurations, actor-oriented change, and identification or uncertainty. DCR belongs under identification and uncertainty on Day 14 because it recalculates dependence-aware uncertainty for a declared dyadic regression. The final row names the frontier tools without pretending they replace the course methods. Ask where each student's project sits and which neighboring cell contains the most serious alternative.

SLIDE 57: AFTER THE MODELS RUN, WRITE THE SUBSTANTIVE PARAGRAPH

What students see

1. State the substantive question and estimand.

2. Explain why the model represents the dependence that matters.
3. Give the substantive result, not a coefficient tour.
4. Name the diagnostic, limitation, and result that would change the choice.

“ANF and ASIM cooperate with far more organizations than the rest. In the initial ERGM, ties between organizations with shared partners receive a clear positive conditional association. After every organization receives its own baseline propensity to cooperate, the shared-partner interval includes zero and the simulations cover the two selected scalar yardsticks. This saturated diagnostic does not clearly separate an additional shared-partner pattern from the hub structure. It is not evidence that closure never occurs.”

Model choice is an argument with a flip condition, not a ranking. Advanced methods should sharpen the substantive claim and its limits.

What to emphasize

Finish with one paragraph: question, model choice, substantive pattern, diagnostic, limit, and the observation that would move the project to the runner-up. Be explicit that the final ERGM uses 29 organization-specific sociality parameters for 85 ties and that its simulations assess two selected scalar yardsticks, not every possible degree and transitivity feature.

SLIDE 58: WHAT TO TAKE WITH YOU

What students see

1. Define actors, ties, eligibility, missingness, and time before modeling.
2. Identify where the important variation and dependence live.
3. Match the model to the estimand rather than to a preferred software package.
4. Diagnose the feature that would change the substantive conclusion.
5. Treat latent quantities, graph statistics, temporal mechanisms, and robust standard errors as different quantities that require different interpretations.

The best network model is not the one that explains everything. It is the one whose question, assumptions, diagnostics, and claim you can state clearly and defend.

What to emphasize

Closing: projects, not another package. One-paragraph model-choice argument plus the diagnostic that would make it change.