

MASSEY RATINGS

- **Inputs**

The computer considers the outcome of each game. Optionally, some games may be assigned more or less weight than usual (e.g. playoffs or exhibitions). All games are analyzed in concert, so that the impressiveness of each win is constantly being re-evaluated in light of the other results.

- **Model**

The mathematical rating model is based on maximizing the retrospective probability of the observed game results. Each outcome generates a “force” on the teams involved, attempting to push the winner above the loser. But as the algorithm adjusts the ratings to create an equilibrium, a ripple effect occurs throughout the entire network of teams. Opponents’ opponents’ opponents’ ... ad infinitum are influenced by some degree. By chains of interaction, it is possible to compare teams that are geographically dispersed.

- **Margin of Victory**

Each game score is translated to a probability that the winner is really the better team. A narrow win might translate to 58%, while a blowout of gives 98%. The cap of 100% enforces diminishing returns to running up the score.

- **Pace**

The model does not discriminate against styles of play that result in fewer total points. In football for example, a team that typically wins games 24-10 may be considered more impressive than a team that wins 49-31 barnburners.

- **Strength of Schedule**

Each team is measured by its performance relative to the opposition faced. Ratings are strength of schedule are calculated simultaneously so that schematically **rating = performance + strength of schedule**. The model is able to accurately compare e.g. a team that went 9-1 against weak competition to a team that went 6-4 against a brutal schedule.

- **Mismatches**

The model drives most of its information from games between teams of similar strength; therefore there should be no incentive to scheduling inferior opponents, since there is limited reward to wins and potentially large downside in the unlikely event of an upset. Regarding strength of schedule, it is more difficult for an elite team to face #2 and #100 than to face #39 and #40.

- **Head-to-Head**

Sometimes lower ranked teams defeat higher ranked teams. These “upsets” are inevitable and should be tolerated since each team is rated according to its entire “body of work”. A single head-to-head result is not always consistent with rankings derived as a best fit for the entire season.

- **Objectivity**

All teams are treated equally and anonymously, without regard to name brands or affiliations. The computer can assess bad or mediocre teams just as well as the teams at the top.

- **Rating and Power**

A team’s rating is designed to reward the most impressive resumes, giving more credit to wins, regardless of how dominant they were. In contrast, the “Power” of a team is more indicative of the true strength of the team. The power of a team is broken in to offensive and defensive components, which can be combined to forecast typical scores for a given matchup, as well as the associated probabilities of winning.