



Ipsos Poll Conducted for Reuters

College Athletes 04.22.2015

These are findings from an Ipsos poll conducted for Thomson Reuters April 21-27, 2015. For the survey, a sample of 2,862 Americans, ages 18+ were interviewed online. The precision of the Reuters/Ipsos online polls is measured using a [credibility interval](#). In this case, the poll has a credibility interval of plus or minus 2.1 percentage points. For more information about credibility intervals, please see the appendix.

The data were weighted to the U.S. current population data by gender, age, education, and ethnicity. Statistical margins of error are not applicable to online polls. All sample surveys and polls may be subject to other sources of error, including, but not limited to coverage error and measurement error. Figures marked by an asterisk (*) indicate a percentage value of greater than zero but less than one half of one per cent. Where figures do not sum to 100, this is due to the effects of rounding. To see more information on this and other Reuters/Ipsos polls, please visit <http://polling.reuters.com/>.

COLLEGE ATHLETES

Q1. Please indicate how much you agree or disagree with the following statements:

	<u>Strongly agree</u>	<u>Somewhat agree</u>	<u>Somewhat disagree</u>	<u>Strongly disagree</u>	<u>Don't know</u>	<u>TOTAL AGREE</u>	<u>TOTAL DISAGREE</u>
I am a fan of college sports and athletics	26%	26%	15%	26%	6%	52%	41%
College athletes should receive a small salary from the athletic program to pay for living expenses	20%	31%	13%	22%	13%	51%	36%
College athletes should be able to profit from autographs, appearances or being featured in things like video games	19%	30%	14%	24%	13%	49%	38%
College athletes in Division 1 schools should be paid a comparable salary to minor league athletes	11%	20%	19%	33%	17%	31%	52%
College athletes should be able to negotiate a salary in much the same way professional athletes do already	10%	20%	21%	35%	14%	29%	57%

How to Calculate Bayesian Credibility Intervals

The calculation of credibility intervals assumes that Y has a binomial distribution conditioned on the parameter θ , i.e., $Y|\theta \sim \text{Bin}(n, \theta)$, where n is the size of our sample. In this setting, Y counts the number of “yes”, or “1”, observed in the sample, so that the sample mean ($\bar{y}$) is a natural estimate of the true population proportion θ . This model is often called the likelihood function, and it is a standard concept in both the Bayesian and the Classical framework. The Bayesian ¹ statistics combines both the prior distribution and the likelihood function to create a posterior distribution. The posterior distribution represents our opinion about which are the plausible values for θ adjusted after observing the sample data. In reality, the posterior distribution is one’s knowledge base updated using the latest survey information. For the prior and likelihood functions specified here, the posterior distribution is also a beta distribution ($\pi(\theta/y) \sim \beta(y+a, n-y+b)$), but with updated hyper-parameters.

Our credibility interval for ϑ is based on this posterior distribution. As mentioned above, these intervals represent our belief about which are the most plausible values for ϑ given our updated knowledge base. There are different ways to calculate these intervals based on $\pi(\theta/y)$. Since we want only one measure of precision for all variables in the survey, analogous to what is done within the Classical framework, we will compute the largest possible credibility interval for any observed sample. The worst case occurs when we assume that $a=1$ and $b=1$ and $y=n/2$. Using a simple approximation of the posterior by the normal distribution, the 95% credibility interval is given by, approximately:

$$\bar{y} \pm \frac{1}{\sqrt{n}}$$

For this poll, the Bayesian Credibility Interval was adjusted using standard weighting design effect $1+L=1.3$ to account for complex weighting²

Examples of credibility intervals for different base sizes are below. Ipsos does not publish data for base sizes (sample sizes) below 100.

Sample size	Credibility intervals
2,000	2.5
1,500	2.9
1,000	3.5
750	4.1
500	5.0
350	6.0
200	7.9
100	11.2

¹ *Bayesian Data Analysis, Second Edition, Andrew Gelman, John B. Carlin, Hal S. Stern, Donald B. Rubin, Chapman & Hall/CRC | ISBN: 158488388X | 2003*

² Kish, L. (1992). *Weighting for unequal Pi*. *Journal of Official, Statistics*, 8, 2, 183200.