

Competition Act Rules on Comparative Advertising Clarified by Recent Court Decision

Date: November 26, 2013

Lawyers You Should Know: H. Todd Greenbloom

Original Newsletter(s) this article was published in: Blaneys on Business: December 2013

Businesses sometimes advertise that their products and services are bigger, faster and better than their leading competitors'.

How far can they go with these comparative claims before they run afoul of the false and misleading advertising provisions of Canada's *Competition Act*?

The Ontario Superior Court has shed fresh light on this question. It illustrates that before making any advertising claims, especially comparative claims, the advertiser must:

- determine who the target might be;
- understand how the ad will be perceived, and
- ensure that the testing used to verify the ad's claims is recognized by the market research industry generally and is used commonly in the applicable industry.

While advertising can be expressive, it is not a Wild West show where anything goes. There are various sets of rules that govern comparative advertising. One set deals with trademark law (and for that reason comparisons are often made with unnamed "leading competitors"). Another set of rules emerges from the *Competition Act*, and it is with those rules that this article deals.

The *Competition Act* prohibits a person from making false or misleading claims and also from making a representation regarding "the performance, efficacy or length of life of a product that is not based on an adequate and proper test thereof."

The consequences of making a false or misleading claim are significant -- imprisonment for up to 15 years and, for corporations, a fine of up to \$10 million for the first offence. Inadequate testing of claims that are made can lead to fines of up to \$10 million for a corporation.

Some of the questions then become: What is adequate and proper testing? How exhaustive does the testing have to be? Under what conditions should the test be conducted? If there are several tests, which should be chosen? How large should sample sizes be? The recent case of *Canada (Competition Bureau) v. Chatr Wireless Inc.* might provide some answers.

Chatr is a brand that Rogers Communications established specifically 1) to compete with new wireless carriers in the prepaid zone/unlimited text and talk segment of the wireless industry and 2) to avoid losing significant market share, as had happened to incumbent carriers in the U.S who waited too long to compete for this segment after it emerged. Rogers determined that price was not going to differentiate it from its competitors. Rather, it believed that it was going to derive its competitive advantage from the quality of its service compared to such competitors as Wind Mobile, Public Mobile and Mobilicity.

Rogers had a more mature network and a frequency that penetrated indoors better than its competitors'. In addition, and of considerable significance, when a customer left Chatr's own zone, it was transferred seamlessly to Rogers' main network. By comparison, the other cellphone companies' customers would be disconnected on switching zones and would have to reconnect to continue a conversation.

In order to capitalize on its advantages, Rogers made the following two advertising claims: "Fewer dropped calls than new wireless carriers" and "no worries about dropped calls."

The other carriers initiated a *Competition Act* complaint about the advertising slogans of Rogers. As support for the claims being false and misleading, the other carriers relied on "switch tests" to look at the data produced by the switches that directed calls on the various networks. Rogers, on the other hand, relied on "drive tests" where calls were made from phones from the applicable carriers in cars to fixed landlines in the same location. The car drove around the fixed route and the calls were monitored for various performance criteria over the course of the drive. Needless to say, the different tests produced different results.

To a large extent, the Chatr case was determined on the basis of how tests should be conducted to verify claims. As background, it should be noted that the Advertising Standards Canada (ASC) Guidelines provide that "Research to support a specific comparative claim against another product or service should follow published standards of the market research industry, or generally accepted industry practices" and "The assessment of comparative advertising research should be based on two principles: validity and reliability."

In the end, the judge, Mr. Justice Frank N. Marrocco, Associate Chief Justice of the Ontario Superior Court, found more favour with the drive test than the switch test. It was determined that "benchmark drive testing is accepted universally as a way of comparing key performance indicators, including dropped call rates, on different networks. Drive testing does not have to be a perfect test to be an adequate and proper test." On the other hand, it was determined that the switch test, which presumably accounts for most calls, was not an appropriate test, since the way the information is collected over different systems is not the same. In other words, apples

are not being compared to apples, but rather to oranges. The drive test was accepted because it was used universally to measure performance.

Other significant findings include the following:

- The proper consumer perspective to be applied to the ads in question was not any credulous and inexperienced consumer but rather a credulous and *technically inexperienced* consumer of wireless services. The target for the ads had some experience because they knew they wanted unlimited talk and text.
- The claim would be perceived to be true for each city where the services were offered vs. an average over all the cities.
- Despite the ASC Guidelines provision that comparative performance claims should not be made when the difference is barely discernible to consumers, in this case “distinguishable” was the key vs. “discernible.” This conclusion was reached because “every dropped call matters” and a “credulous and technically inexperienced consumer would choose a network that offered fewer dropped calls to avoid the possibility of an important call being dropped.”
- For the reasons set out above, the difference did not have to be different statistically.
- Even though Rogers had a better network, and the logical inference is that a better network leads to fewer dropped calls, the testing still had to be conducted.
- The testing does not have to meet the standards of an academic paper.
- The test has to be done *before* making the claim. A test done after the claim is not enough, even if it supports the claim.
- The tests do not have to be validated by an independent third party.

As we indicated earlier, the primary lessons learned from the Chatr case are that before making any claims, especially in the context of a comparative claim, the advertiser must know who the target might be, understand how the ad will be perceived, and must make sure that the testing used is recognized by the market research industry generally and is commonly implemented in the applicable industry.

Finally, in the Chatr case, Rogers was helped by the large testing infrastructure that had been built to conduct drive testing in the wireless industry.